

Born-Digital Humor and Context Saturation: Using Generative AI to Preserve and Interpret Russo-Ukrainian War Memes

Anna Rakityanskaya, Slavic Librarian, Harvard Library, Humanities and Social Sciences Collections and Services (HSSCS), Harvard University, rakityan@fas.harvard.edu

Memes of the full-scale invasion phase of the Russo-Ukrainian War (2022-) are characterized by their abundance, varied focus on multiple aspects of the war, and high ephemerality due both to their physical container and their complex multimodality and intertextuality. Understanding the humor of memes is vital for preserving the memes' integrity as communication units, because their meaning often cannot be gleaned from the explicit properties of the meme and emerges only through an understanding of the implicit message.

This paper focuses on finding a solution for preserving the semantic integrity of memes in digital archiving by testing the application of generative AI in their interpretation. Using the framework of Suls' (1972) incongruity resolution theory and Yus's (2017, 2021) system of incongruity patterns, the study tests the ability of ChatGPT 4o to process the content of memes to produce metadata with an accurate explanation of their humor. The author introduces the concept of the meme context saturation spectrum as a tool for evaluating the correlation between a chatbot's performance and the amount of contextual information embedded in a meme. The study also measures the chatbot's success in processing both current and retrospective material.



Introduction

The phenomenon of Russo-Ukrainian War memes has attracted the attention of media, researchers, and digital archivists since the beginning of the full-scale invasion in February 2022 (Rakityanskaya, 2025). These memes are characterized by their abundance, varied focus on multiple aspects of the war, and high degree of ephemerality. Because of the sheer number of memes and the fact that they react to rapidly changing developments both in the war and the public discourse about it, their meaning quickly fades. Moreover, the specific cultural context (pivotal to an understanding of Russo-Ukrainian War memes) is quite complex, multi-layered, and sometimes exclusive, which makes some of its components inaccessible even to their immediate consumers. An understanding of the humor of these memes is crucial for more than just entertainment purposes. More importantly, it is a condition *sine qua non* for preserving the meme's integrity as a communication unit in which the meaning emerges as a result of the processing of both explicit properties (image and text) and the implicit message, often rooted in contextual references. In this study, we focus on finding a solution for preserving the meaning of memes from the standpoint of digital archiving by testing the use of a generative artificial intelligence chatbot for interpreting them in correlation with the chatbot's ability to process contextual information.

Theoretical Considerations

The term 'meme' was coined by Richard Dawkins in 1976 in his application of gene theory to the evolution of human culture. Dawkins (1976: 206) used the term to describe 'a unit of cultural transmission,' that propagates itself 'by leaping from brain to brain' through imitation. L. Shifman (2014: 41) developed the concept of internet memes, pointing out their dynamic nature as 'a group of digital items' that are shared and updated by many Internet users.

When discussing the communicative aspect of memes, it is crucial to acknowledge that, as any genre of (multimodal) humorous text, they assume a prepared audience familiar with the many layers of relevant context (Wiggins, 2019; Milner, 2016). Many scholars underscore the essential role of cultural background in the perception of any images (Boylan, 2020) and even more so, for understanding political memes (Anderau and Barbarrusa, 2024; Lynch, 2022; Kirner-Ludwig, 2022).

When it comes to memes that have emerged around a specific event in a non-English language environment, comprehending them requires an even more complex form of background knowledge, as they combine two types of context: the general context inherent to international meme-making practices and the local context rooted

in the regional culture. A. Denisova (2019: 54–68) points out the need to understand the local media context, which includes contemporary history, national identity, political environment, media environment, and media audience. L. Laineste and P. Voolaid (2016: 26) in their case study of Estonian memes dissect their intertextuality, which assumes the local ‘cultural memory’ of the Estonian-language community (including Soviet-era cultural practices), as well as familiarity with Russian- and English-language culture.

The situation is further complicated by the exclusive nature of online humor in general. In the multidimensional world of information, understanding memes is dependent not only on knowledge of factual content and familiarity with both general and regional cultural references but also on access to subcultural references specific to certain social networks and online communities. These communities create and share jokes and memes that are understood within the group but impenetrable to outsiders (Milner, 2016; Phillips, 2015; Gal, Shifman, and Kampf, 2016).

The notion of context becomes critical when viewed in relation to the problem of the memes’ ‘afterlife’ (a term coined by Nahon and Hemsley (2013: 124–142) to describe a stage in the existence of any viral piece of internet communication when it is forgotten and no longer communicates information). The digital ‘afterlife’ of memes occurs when all or most of the semiotic strings that link the meme to social, political, and technological context disintegrate due to their loss of relevance and the selectiveness of human memory. No longer linked to its context, a meme ceases to exist as a multimodal and multitextual unit, loses its humor, and becomes a meaningless digital image.

The task of archival preservation of internet memes entails extracting them from the original semi-closed online environments (in which memes might be available physically, but not necessarily intellectually) and making them intellectually accessible to virtually anyone. This means that, in addition to producing regular metadata describing the visible aspects of a meme and its general subject, archivists must also preserve its meaning and unique humor, a daunting task that goes beyond traditional archiving practice. In the words of García López and Martínez Cardama (2020: 896), ‘memes cannot be preserved as isolated digital objects devoid of the context that affords them full meaning.’ This task is well understood by the administrators of *Know Your Meme*, a ‘cyber warehouse’ [Serb. сајбер складиште] of mostly English-language memes, as pointed out by Knežević (2023).

Memes that form part of public discourse during a specific event (like the Russo-Ukrainian War) are important historical documents, and they must be preserved as primary sources. Details of day-to-day combat and civilian experience, viral words and concepts fill social media in some cases only for a few days and then disappear from

collective memory, but remain encoded in memes. However, without their context, these memes can appear completely opaque (as demonstrated by Rossi and Bondarenko, 2024) or be interpreted in a way that is directly opposite to their original meaning, even by those internet users who share the political views of the meme creators. One example is the meme in **Figure 1**. Its earliest known version was published on Reddit on February 27, 2022 (Busdriver242, 2022). The meme references a video published as early as February 25, 2022 (the second day of the full-scale invasion). In the video (*The Guardian*, 2022) a Ukrainian woman tells Russian occupant soldiers to put sunflower seeds in their pockets, ‘so at least sunflowers will grow when [your bodies] are lying here’ (translation my own). Thereafter, in the semiotic system of Ukrainian memes, an image of a field full of sunflowers has stood for ‘dead Russian invaders.’ The memetic concept of sunflowers growing from the graves of Russian invaders is an example of context that is essential for understanding and properly categorizing the meme. In the absence of that context, the meme is occasionally interpreted in a directly opposite way, as supportive of the Russian propaganda narrative, suggesting that Russian troops bring peace to Ukraine: they enter bearing arms and leave it flourishing.¹

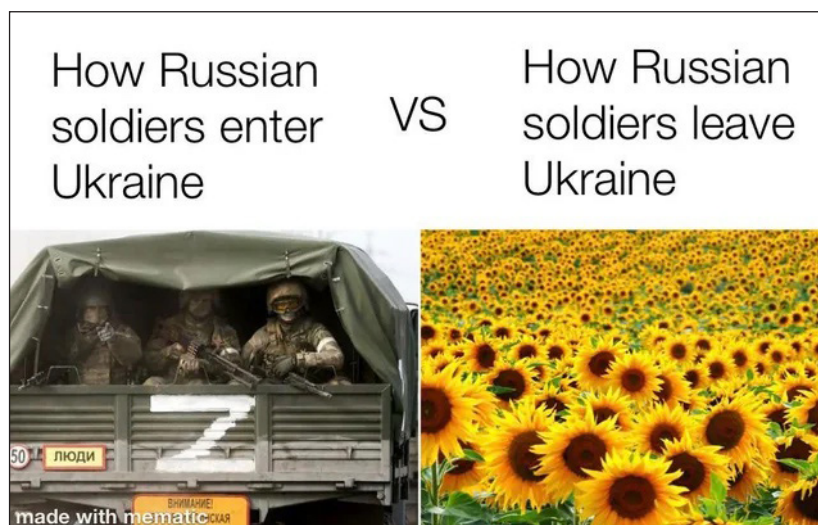


Figure 1: How Russian soldiers enter Ukraine vs. How Russian soldiers leave Ukraine. Posted by Busdriver242 on *Reddit*, February 27, 2022. https://web.archive.org/web/20250309135438/https://www.reddit.com/r/memes/comments/t2tyvs/it_doesnt_look_good_for_them/ [Last Accessed 31 October 2025].

Given the vital importance of preserving the meaning of memes in meme archiving, particularly in light of the vast number of memes produced around the Russo-Ukrainian War, we are particularly interested in exploring the potential of artificial intelligence

¹ See such decontextualized interpretation in Chen et al. (2023: 30).

(AI) for processing them. Fortunately, in recent years, many researchers have turned their attention to the application of AI in the humanities and social studies (Thapa et al., 2025). Significant progress has been made in research on using AI to recognize stance (particularly hate: Alam et al., 2024; Hossain et al., 2024; Karim et al., 2023), offensive content (Sharma et al., 2020), humor (Wang, 2023) and its specific forms (like sarcasm: Kumari et al., 2024; Cai, Cai, and Wan, 2019) in social media posts and memes. Of particular interest are studies conducted on non-English material, like Telugu (Bellamkonda, Lohakare, and Patel, 2022) or Spanish (Chiruzzo et al., 2021). Equally relevant are the studies of the ability of different AI tools to recognize objects in images and to produce metadata for the images (Thammastitkul, 2023). L. Soriano-Gonzalez and J. Belda-Medina (2024) undertook an impressive experiment exploring the use of different Large Language Models (LLMs) to interpret memes from the standpoint of pragmatics, which involved not only processing the explicit elements of the memes, but also understanding their implicit references.

While scholars overall find the application of AI in the interpretation of memes a beneficial and promising technology, many note its lack of cultural context as a significant drawback (Soriano-Gonzalez and Belda-Medina, 2024; Wang, 2023). V. Prabhakaran, R. Qadri, and B. Hutchinson (2022) point out ‘cultural incongruencies’ resulting from AI systems being developed in a small number of countries and trained on data limited to certain dominant cultures. Our study contributes to the diversification of research into the applications of AI in the humanities, as it is concerned with memes produced in connection with a specific contemporary historical event (the Russo-Ukrainian War), taking place in a specific geographical area (Ukraine and neighboring countries).

This study employs as its theoretical framework J. Suls’ incongruity resolution (IR) theory and F. Yus’s incongruity patterns taxonomy. Suls’ (1972: 81–100) IR theory describes the process of humor appreciation in verbal jokes and cartoons as a sequence of two stages: 1. the ‘joke recipient’ encounters an incongruity (which in the case of a cartoon means that the caption ‘disconfirms’ the image in some way); 2. the recipient finds a ‘cognitive rule’ that reconciles the incongruity. F. Yus (2017: 105) further developed the IR theory by discovering two basic patterns of incongruity in jokes: frame-based (when a joke clashes with ‘the hearer’s construction of an appropriate mental situation (frame, schema, script, etc.)’) and discourse-centered (when a semantic conflict occurs in processing ‘the verbal content of the joke’, producing word-play jokes). Later, in his study of meme-specific humor perception, Yus (2021: 137) introduced a third type of incongruity, discourse-image based, which ‘involves an inferential clash when inferring meanings that need a convergence between the partial meanings of the text and the image’.

We propose a fourth type of incongruity in memes: *frame-image based incongruity*, where the image itself creates an inferential clash with the perceived mental frame. Yus did not use this term, although he described an IR pattern in memes where an image plays an essential role in combination with frame-based incongruity. Not surprisingly, the researcher found a negligible percentage of such memes in his sample, which contained only memes constructed as a combination of image and text. Memes that have no text were not included in the sample.

The meme in **Figure 2** is an example of a frame-image incongruity meme. The meme is composed of two photos featuring a headshot of a man and a meme macro ‘monkey with lips’ (Meme-arsenal, 2025). In the total absence of captions, however, the meme lends itself to a certain level of comic interpretation even without context, because it offers an obvious visual incongruity: a man reacting to having a conversation with a monkey. On this very basic level, the resolution comes from the idea that a man and a monkey cannot have a meaningful conversation. However, this meme was created in the context of the Russo-Ukrainian War, and its actual meaning is more specific. The meme is a response to the interview with Vladimir Putin conducted by the American journalist Tucker Carlson that aired online on February 8, 2024. The photograph of the man is a still of Carlson during this interview. Putin, however, contrary to the viewer’s expectations, is replaced by a monkey, which creates the meme’s incongruity that is resolved by mentally equating Putin’s pompous rants with the unintelligible vocalizations of a monkey. Thus, this frame-image incongruity meme requires knowledge of the context to fully resolve the incongruity created by the combination of the image and the frame.



Figure 2: [Tucker Carlson and a Chimpanzee with lips]. Posted by anonymous in comment thread on NEXTA Live, February 13, 2024. https://t.me/nexta_live/71029 [Last Accessed 13 February 2024].

Context saturation spectrum

For any audience for memes created around a specific historical event (like the Russo-Ukrainian War) in a predominantly non-English information environment, possessing complex contextual information directly affects the ability to retrieve the ‘cognitive rule’ that would allow one to reconcile the meme’s incongruity. A related image, a viral word or phrase, or a viral idea captured in a meme may reference specific events, people, or cultural practices that are part of the background required to both create and understand the meme. To evaluate the ability of generative AI to ‘read’ contextual information in memes, we developed a meme context saturation spectrum that arranges amounts of contextual information by template type.

While the common understanding of a meme template is an image macro,² we find that memes of the Russo-Ukrainian War have a more complex template structure. If we recognize a template as a stable element of a meme that is subjected to alteration as it is replicated, we must acknowledge that in addition to visual templates, there are also verbal and conceptual templates.

Visual templates. We distinguish between non-contextual and contextual visual meme templates used to build memes of the Russo-Ukrainian War:

1. Non-contextual templates. These images are not linked to the war or do not contain any region-specific visuals; instead, they use universal internet meme templates (image macros) and techniques. These may include:
 - General image macros (e. g., Distracted boyfriend, Peter Parker glasses, Buff Doge vs. Cheems, etc.);
 - Image macros based on classical Western films and TV shows (e. g. *Harry Potter*, *Lord of the Rings*, *Game of Thrones*, *The Simpsons*, etc.);
 - Stock images from broad categories of artwork, computer game screenshots, historical photos, animal photos, etc.;
 - Memes ‘mimicking’ other genres of internet communication (like infographics or screenshots).
2. Contextual images. This group of images is linked to the broadly understood regional context related to the countries engaged in the war. They include:
 - Images drawn from the war coverage by news media and bloggers (e. g. photographs depicting Angelina Jolie and a teenager inside a cafe in Lviv,

² See, for example, Yus’s description: ‘A particularly abundant type of meme at present is the *image macro*, a text-image multimodal discourse made up of one or two text lines at the top and/or the bottom of the meme complemented by an image in the middle with several possible interpretive combinations’ (Yus 2021: 135).

Emmanuel Macron hugging Volodymyr Zelensky, handcuffed Viktor Medvedchuk, a bird's-eye view of bodies of Russian soldiers in a snowy field, a crying Russian woman in a car fleeing from Crimea, Yevgeny Prigozhin shouting, etc.). Some of these images become image macros and generate many memes, while others never achieve memetic status and are used only occasionally.

- Pre-2022 image macros which originated in Ukraine (e.g., Volodymyr Hroisman speaking with a man and the Poltava arsonist);
- Pre-2022 image macros which originated in the wider post-Soviet internet space (Get up Natasha!, And what about goblets?, Monkey Putin, Vatnik, etc.);
- Image macros based on classical Soviet films (*The Night Before Christmas*, *Heart of a Dog*, etc.).

Verbal templates (memetic words/phrases). Some memes are built around an idiom, a viral word, or a phrase that has risen to prominence in the popular discourse around the war (e. g. 'Chornobaivka', 'Red lines', 'Palianytsia', 'Kyiv in 3 days', 'We haven't really started anything yet', etc.). We view these as textual templates that are continuously replicated and edited with the help of images.

Conceptual templates (memetic ideas/concepts). A two-dimensional meme can be built on a conceptual template if it uses a viral idea or a folkloric story that is part of the popular narrative but has not been encapsulated in a single word or phrase. Some examples of these viral ideas/concepts in memes of the Russo-Ukrainian War are sunflowers growing on graves of Russian soldiers, Kherson watermelons to be tasted after the liberation of Kherson, Ukrainian dogs, foxes and other animals getting fatter as a result of eating enemy corpses, Ukrainian farmers stealing tanks and old ladies downing drones with a jar of pickles, Ukrainian birds and mosquitos fulfilling combat missions, Ukrainians offering eggs instead of rare earth minerals to Americans, Putin's weakness for historical excurses, and many others. Each of these narratives originated at a specific (and often documented) point in time, became widely popular, and was propagated in various modes, including digital images and videos.

It must be acknowledged that in two-dimensional memes, visual templates have a more prominent position, since a visual template is always present, while a verbal or conceptual template is optional. However, when a non-visual template is employed, it is essential to recognize it, especially because when combined with a non-memetic (occasional) image, it becomes the meme's main structural element.

To measure the amount of context in memes, we rated the combinations of visual, verbal, and conceptual templates in the memes by the degree of context saturation

(CS) they contain (**Table 1**). By CS, we understand the perceived amount of contextual information embedded in a meme's structural elements. Three groups of context saturation patterns emerged around three basic types of visual templates: non-contextual image (NCI), contextual image (CI), and a combination of the two (NCI + CI), to which a memetic idea (MI) or a memetic word/phrase (MWP) can be added. Although theoretically one could expect options that would combine MI and MWP, our material did not reveal such instances. We can speculate that the reason MI and MWP do not (or possibly, rarely) coexist in one meme might be the fact that MIs and MWPs describe different concepts (in other words, they are mutually exclusive) and memes present highly focused messages that usually mention one such concept. However, this theory needs further exploration using a larger meme sample.

context saturation pattern	context saturation level
NCI	1
NCI + MI	2
NCI + MWP	3
NCI + CI	4
NCI + CI + MI	5
NCI + CI + MWP	6
CI	7
CI + MI	8
CI + MWP	9

Table 1: Meme context saturation spectrum.

We rated the resulting 9 groups from 1 (the least contextual NCI), to 9 (CI + MWP, a combination that assumes the highest level of precise background information necessary for deciphering the meaning).

According to our initial hypothesis, ChatGPT would be particularly challenged by the context-heavy high (7–9) and medium (4–6) segments of the CS spectrum.

Our study also had a diachronic aspect: we were interested in how well AI would analyze the material from 2022 through 2025 retrospectively. Evaluating the ability of ChatGPT to analyze retrospective information accurately is a very important aspect of testing because retaining information over time is pivotal for the archival processing of memes in their afterlife. We hypothesized that ChatGPT acquires information gradually, retains it, and performs equally well analyzing both older and newer material.

The Study

The study was conducted using ChatGPT 4o, a generative AI chatbot with a proven capability for understanding humor (Soriano-Gonzalez and Belda-Medina, 2024) and processing Ukrainian language texts (Syromiatnikov, Ruvinskaya, and Troynina, 2024). We selected a sample of 120 memes from the SUCHO Meme Wall dataset.³ The guiding principle for the selection was a balanced representation of basic structural elements and publication dates. The resulting sample included 58 memes built on contextual images and 62 using non-contextual ones. Forty memes were published in 2022, 35 in 2023, 37 in 2024, and 8 in 2025 (the number reflects the share of 2025 memes at the time of the testing in the spring of that year).

For each meme, the researcher manually assigned an incongruity pattern and a CS level.⁴ The chatbot was introduced to the theory of incongruity patterns through the text of *Incongruity-Resolution Humorous Strategies in Image Macro Memes* (Yus, 2021) and the researcher's own presentation of frame-image-based incongruity. The chatbot was then subjected to a series of tests aimed at evaluating its performance in two major testing categories:

1. Understanding and explaining humor. This category included two tests:
 - Evaluating the ability to recognize incongruity patterns. ChatGPT was asked to assign an incongruity pattern tag to each meme. The results were then compared to the ones assigned by the researcher, and the accuracy of the matches was assessed.
 - Evaluating the ability to produce a coherent narrative explanation of the comic effect. ChatGPT was given a task to write a short (under 150 words) narrative text explaining the comic effect of a meme from the point of view of IR theory. The researcher then evaluated the accuracy of the text.
2. Evaluating the chatbot's ability to identify specific elements of a meme's structure and content (expressed both explicitly and implicitly) that require access to contextual information:
 - A visual template used;
 - A viral word/phrase included or implied;
 - A viral memetic concept implied;
 - An event or a news item referenced;
 - A real person featured or implied.

³ SUCHO Meme Wall dataset is available as a CSV file: <https://memes.sucho.org/about/>.

⁴ It must be acknowledged that manual coding of memes and assessment of Chatbot's accuracy by the same researcher presents a possibility of a bias and should be considered a limitation of the study.

The results of the chatbot's performance in each of the seven tests were evaluated for accuracy on a scale of 1–3, where 1 was assigned to wrong answers, 2 to partially correct, and 3 to correct answers (including no answer if no data were available in the meme). Average scores in each category were then studied in correlation with CS level and the year the meme was posted.

The following standard prompt was used for the test:

Please analyze this meme related to the Russo-Ukrainian War and provide the following information:

- *The original text;*
- *Its translation into English;⁵*
- *The specific event or news item the meme references;*
- *If the meme features or mentions a real person, please name that person;*
- *The visual template(s) / macro(s) the meme uses;*
- *If the meme uses a viral word or phrase related to the Russo-Ukrainian War, please give it in the original language and in English translation;*
- *If a meme uses a popular (memetic) idea or concept related to the Russo-Ukrainian War, please name it;*
- *The comic effect in less than 150 words from the POW of incongruity resolution theory;*
- *Is the incongruity frame, discourse, image-frame, or discourse-image based? Please code the incongruity patterns as FB, DB, FIB, DIB. Please do not add any other text in this field.*

Findings

Incongruity pattern identification

ChatGPT was asked to assign an incongruity pattern to each meme, based on the following taxonomy:

DB: Discourse-based (the joke is verbal and can be understood without an image)

FB: Frame-based (the joke is based on a perceived situation, i.e., frame)

DIB: Discourse-image based (the joke is based on the clash between the text and the image)

FIB: Frame-image based (the joke is based on the clash between the frame and the image).

⁵ The text and its translation were collected for auxiliary purposes only.

We found that the match between the human- and chatbot-assigned incongruity tags was 38.46% (50 memes), partial match 23.08% (30 memes), and no match 38.46% (50 memes).⁶

The correlation between the accuracy of incongruity pattern identification and the meme's CS level (**Figure 3**), despite the overall upward trend toward the higher saturation end of the spectrum, is very uneven, especially in the middle section of the spectrum, where we find both the highest and the lowest scores for the most complex, tripartite template categories (NCI + CI + MI and NCI + CI + MWP, respectively). The addition of memetic ideas and verbal memes appears to have enhanced the processing of non-contextual-image memes, but hindered it for memes utilizing a contextual image as a visual template.

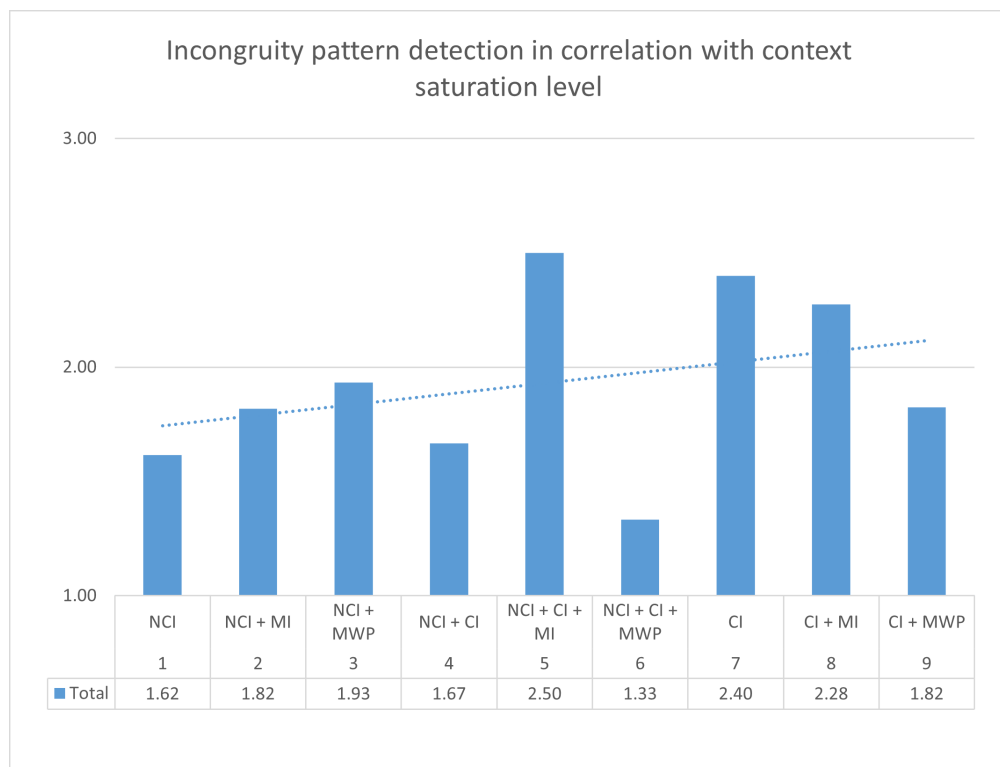


Figure 3: Incongruity pattern detection in correlation with context saturation level.

We compared the accuracy of incongruity tag assignment for memes posted in each of the years of the full-scale war (2022–2025). The results indicate an uneven level of accuracy, with a declining trend for the most recent memes (Figure 4).

⁶ All calculations and graphs were produced in Excel.

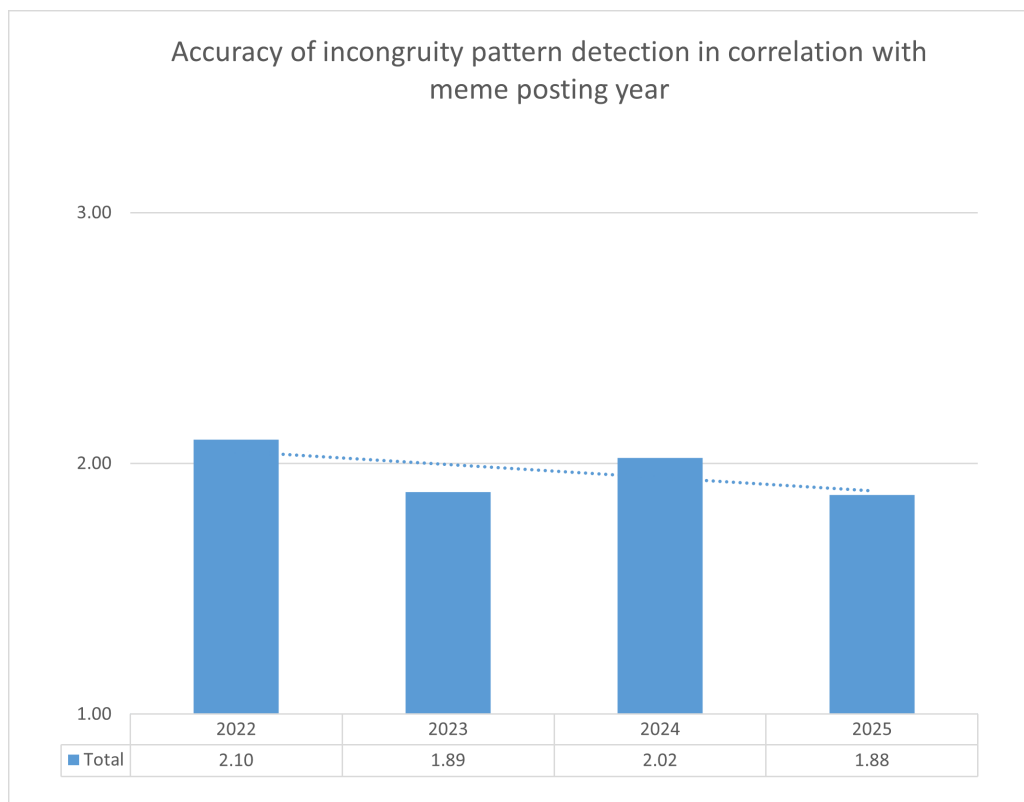


Figure 4: Accuracy of incongruity pattern detection in correlation with the meme posting year.

Comic effect narrative from the point of view of the IR theory

We asked ChatGPT to produce a narrative explaining each meme's comic effect, based on the IR theory, in 150 words or fewer. We view such a narrative as a prototype for a descriptive metadata element of meme archiving. 50% of responses were rated as adequate, while 30.77% were rated as completely incorrect and 19.23% as partially correct.

Overall, the test confirmed our hypothesis that memes from the lower end of the CS spectrum will be the easiest for the chatbot to understand (**Figure 5**). The highest degree of accuracy was demonstrated in processing memes built on non-contextual images. The lowest score was earned for interpreting memes built on a combination of non-contextual and contextual images. At the same time, the presence of additional contextual elements (a memetic idea or a memetic word or phrase) yielded better results.

In the higher third of the CS spectrum, on the other hand, we observe a performance drop in the presence of a memetic idea, followed by a further decline for memes that combine a contextual image with a verbal meme.

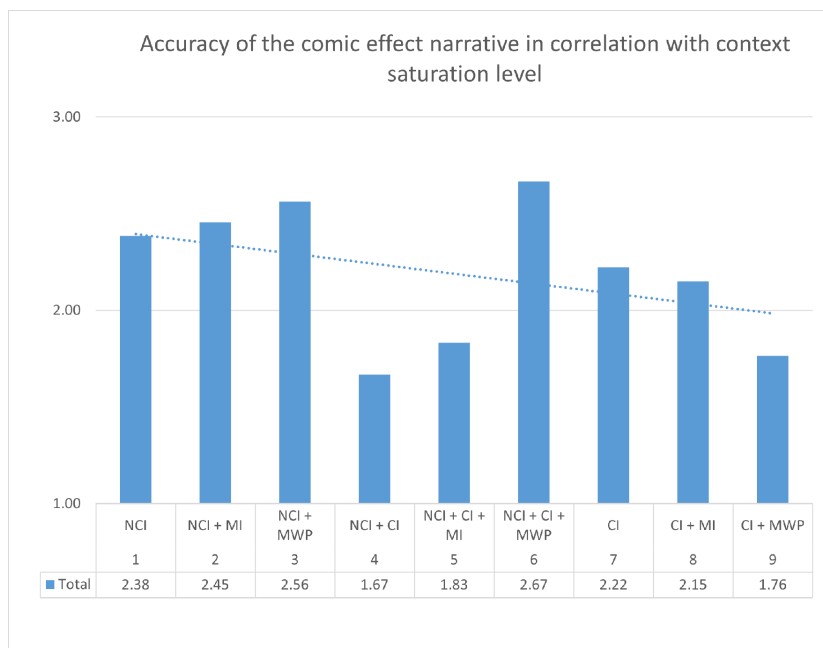


Figure 5: Accuracy of the comic effect narrative in correlation with context saturation level.

The overall score average for narrative explanation of comic effect remains relatively stable over the years, oscillating between 2.16 and 2.25, with a slight dip in 2023 and 2024 (**Figure 6**).

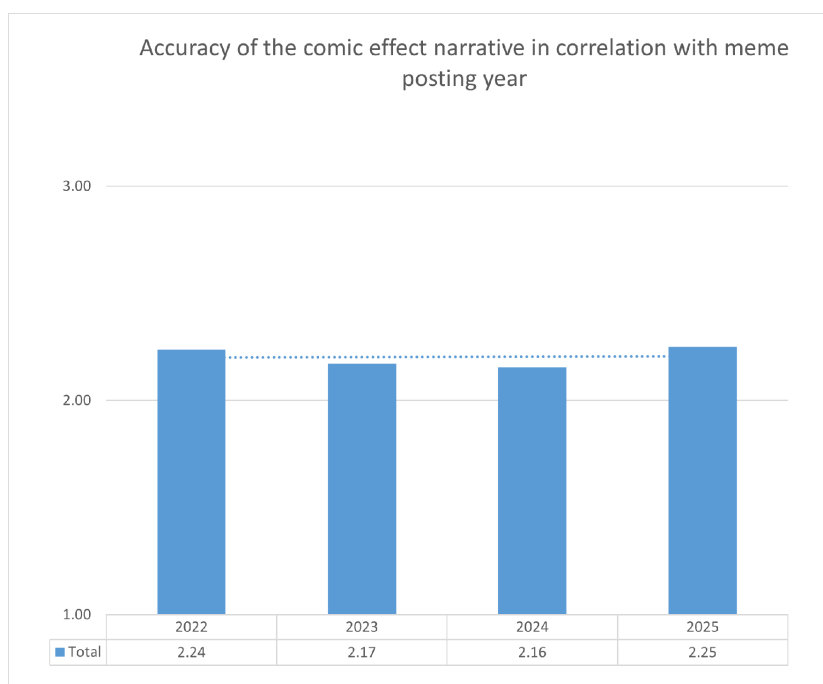


Figure 6: Accuracy of the comic effect narrative in correlation with the meme posting year.

Visual template

ChatGPT correctly identified visual templates (source image types used to create memes) in 75.38% cases, partially correctly in 11.54% and incorrectly in 13.08%.

On the CS spectrum, overall performance declines towards the highest end (**Figure 7**). Identifying visual templates in contextual image-based memes was predictably the most challenging task in this test. A combination of non-contextual and contextual images, on the other hand (the CS spectrum's middle segment), created conditions in which Chatbot showed the strongest performance, with both of the tripartite template categories (NCI + CI + MI and NCI + CI + MWP) consistently rated at 3.

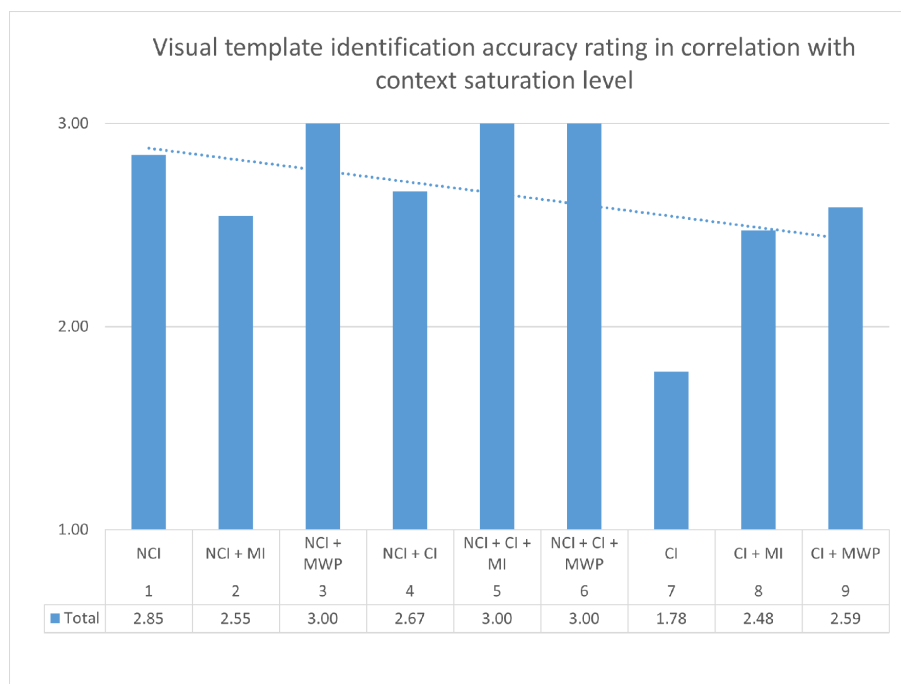


Figure 7: Visual template identification accuracy in correlation with context saturation level.

The templates that ChatGPT was not able to identify included the following categories:

- Internationally lesser-known Ukrainian and post-Soviet image macros, e.g., Flork of Cows; Poltava Arsonist, Get up, Natasha!;
- Viral photos derived from wartime news coverage, e.g., Angelina Jolie and a teenager in a café in Lviv; detained Viktor Medvedchuk;
- Well-known but modified universal image macros; for example, the Distracted Boyfriend meme, where the male character is replaced with a trash bag implying

a fallen mobilized enemy soldier, was identified by ChatGPT as ‘Butterfly metamorphosis meme + visual substitutions’ (elsewhere the same Distracted Boyfriend image macro was identified correctly);

- Certain screenshots from Soviet films; for example, a popular template from the film *Heart of a Dog* (1988) was identified as ‘Retro Russian photo meme’, a still from *Mimino* (1977) was described as ‘Soviet gangster still’.

Partially correct responses had the following issues:

- Lack of precision, e.g., the Father in a bag template (Ukr. Папка в пакеті: Nyzovets, 2023) was identified as ‘Arial message from soldiers shaped with their bodies’, a description which fails to mention that the bodies are in fact dead, i. e., they are corpses;
- Photographs from news coverage featuring politicians were usually described extremely vaguely, e., the template for a meme featuring Donald Trump and Volodymyr Zelensky speaking to each other on the phone, is described as ‘Phone call meme, romantic misfire format’;
- New templates from wartime news coverage photos, e.g. a meme built on a popular distinctive template Chmonya (*Imgflip*, 2025) featuring a diminutive Russian soldier, with Putin’s head photoshopped over Chmonya’s, is described simply as ‘Edited photo of Putin as conscript’;
- Non-war Ukrainian (not known internationally) image macros, e.g. V. Hroisman speaking with a man (*Oboz.ua*, 2017) is described as ‘Ukrainian politician parody meme’;
- A classical image macro with a visual modification, e.g. a meme built on Disaster Girl meme with V. Zelensky’s head photoshopped over the girl’s head, is identified as ‘Fire/explosion background meme’.

All of these failures and shortcomings point to the chatbot’s lack of access to a specific cultural context.

When we compare the Chatbot’s scores identifying visual templates for memes posted in specific years (**Figure 8**), we see a fairly even performance, with the year 2024 standing out slightly.

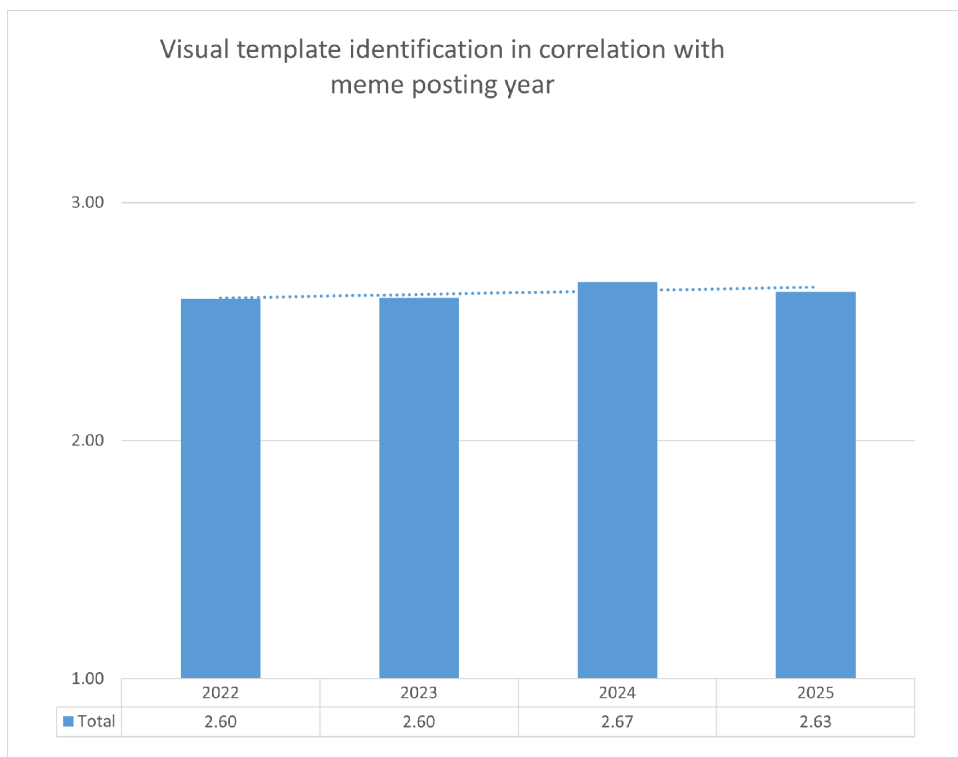


Figure 8: Visual template identification in correlation with the meme posting year.

Viral word/phrase (verbal template)

The task of identifying a viral word or phrase used in a meme proved to be the most difficult for ChatGPT and produced the lowest scores. The success rate was only 33.08%, while 7.69% of the responses were rated as partially correct and 59.23% as incorrect.

We observe an overall upward trend towards the highest-context segment of the CS spectrum (**Figure 9**) and a more balanced performance in that segment. At the same time, there is a pattern of increasing accuracy in the lower and middle segments of the CS spectrum, from an image template alone to an image + verbal template. This pattern suggests that ChatGPT attempted to identify viral words and phrases in memes where none existed (no MWP component). Contrary to its own tactics in identifying a person, when it came to recognizing a viral word/phrase, ChatGPT took very few opportunities to declare an absence of such material, but instead, in majority of cases tried to 'invent' it, coming up with false verbal memes like 'Zaporizhzhia NPP', 'Conscript', 'We'll endure', 'Landing craft', 'Scooter detonation', 'Back and forth', 'Mustache' and 'Worthy place'.

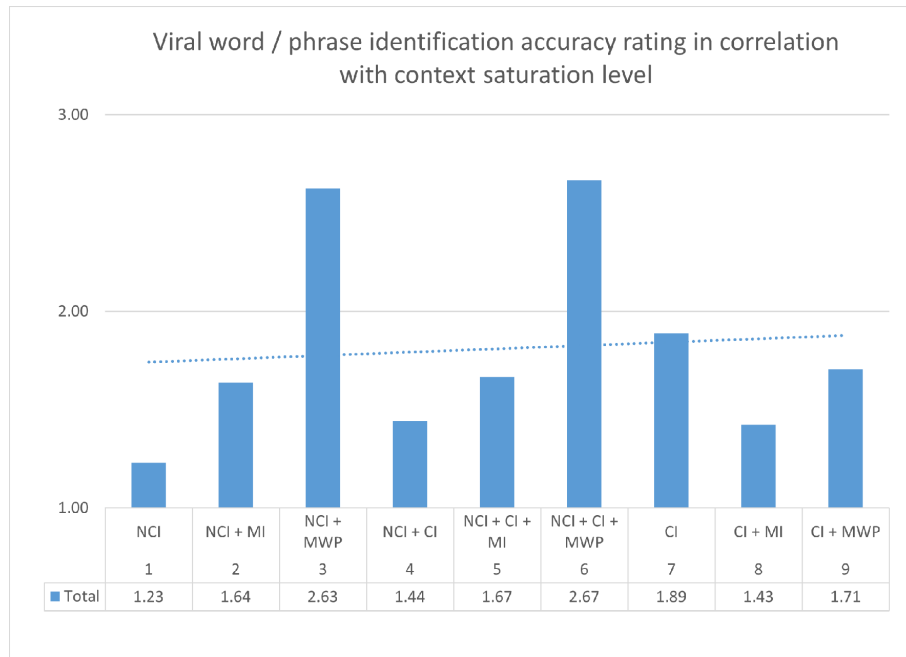


Figure 9: Viral word/phrase identification accuracy rating in correlation with context saturation level.

Memes with the following conditions presented the most difficulty for ChatGPT in this test:

- the actual verbal (explicit) manifestation is not present, instead the word is represented implicitly, through an image (for example, an image of a cotton boll replacing the word ‘cotton’ (Ukr. бавовна), a euphemism for ‘an explosion on Russian or Russia-occupied territory’ in mockery of Russian propaganda language: Vidomenko, 2022);
- a viral word or phrase is modified or clipped. For example, a meme depicting a burning fighter jet going down is captioned as ‘bi bi bi bi bi bi bi’ (**Figure 10**). The Ukrainian creator of the meme used the letters *b* and *i* of the Roman alphabet to render a Russian letter *ы*, absent from the Ukrainian version of the Cyrillic alphabet (a popular workaround). The Russian term ‘khokhly’ (literally: ‘forelocks’) is commonly used as a derogatory nickname for Ukrainians in reference to a traditional Ukrainian male hairstyle. Memes featuring the exaggerated exclamation ‘Khokhlyyy..!’—often ending in a drawn-out ‘yyy!’ (Rus. ЫЫЫ!) or just that ending alone—went viral in January 2024, following a social media post describing a Russian’s frustration with Ukrainians (Kononenko, 2024). These memes satirize Russian propaganda’s tendency to blame Ukrainians for

a wide range of misfortunes. ChatGPT took the caption at ‘face value’ and read it as *bibibibibibi*, interpreting it, in its own comic effect narrative, as ‘childlike sound effects’. Needless to say, the resulting overall interpretation of the meme was grossly incorrect. It is worth mentioning that while analyzing another meme that uses the complete form of this viral word in a caption, ChatGPT identified it correctly as ‘*хохлыыыыыыы* / exaggerated ‘*khokhlyyy*’.



Figure 10: yyyyyyy. Posted by anonymous on NEXTA Live Chat, February 19, 2024. https://t.me/nexta_live/71307 [Last Accessed 19 February 2024].

On a few occasions, ChatGPT truncated the verbal template. In some cases, the result was deemed partially correct, for example: ‘Where were you’ (instead of ‘Where have you been for 8 years?’) is passable. In other cases, such fragmentation led to a completely incorrect interpretation: e.g., ‘The enemy is running after us in shame’ (Rus. Враг с позором бежит за нами), a verbal joke representing Russian propaganda efforts to cover up the retreat of Russian forces was cut to ‘Враг с позором’ and interpreted as ‘shameful enemy’.

Success in identifying verbal templates increased in 2023 and 2024, whereas performance for memes posted in 2022 and 2025 was lower (**Figure 11**).

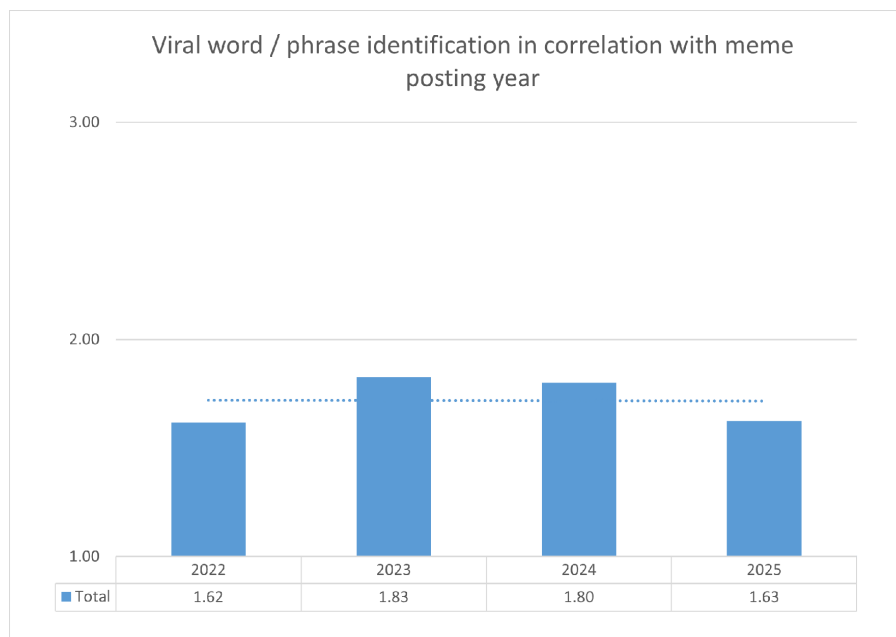


Figure 11: Viral word/phrase identification in correlation with the meme posting year.

Viral idea (conceptual template)

In this test ChatGPT was asked to identify a viral (memetic) idea or concept mentioned or implied in a meme. This is another low-success category, with only 22.31% correct answers, 37.69% partially correct and 40% completely incorrect ones.

As expected, ChatGPT received higher average scores for analyzing memes on the lower end of the CS spectrum, with a stable score growth from level 1 to 3 and a fairly level performance in the high-context segment of levels 7–9 (**Figure 12**). In the middle range of the spectrum on the other hand, amid low scores for levels 4 and 5, we see a sudden uptick in performance for the NCI + CI + MWP category (level 5). In all three segments of the CS spectrum, we observe that memes built exclusively on visual templates are the hardest to analyze. In contrast, memes that combine both a visual template and a conceptual or verbal one consistently earn the chatbot higher performance scores.

A notable flaw in ChatGPT's responses in this category was their vagueness. Generalizations like 'Geopolitical absurdity,' 'Totalitarian equivalency,' or 'Meme urgency culture' seem to reveal the chatbot's response strategy of targeting the broadest possible conceptual area, hoping the correct response will fall within it. Some of these responses were outright incorrect, while others were partially accurate. The partially correct responses belonged to the right general subject area but lacked the necessary focus.

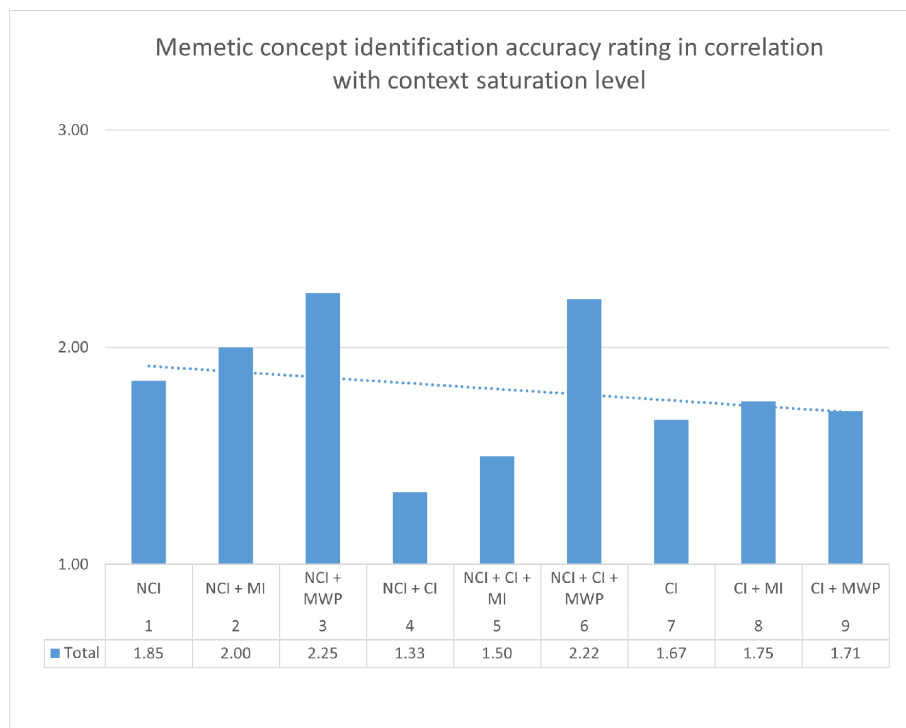


Figure 12: Memetic concept identification accuracy in correlation with context saturation level.

Diachronically (**Figure 13**), the conceptual template was identified with significantly more accuracy for memes published in 2023 than in any other year.

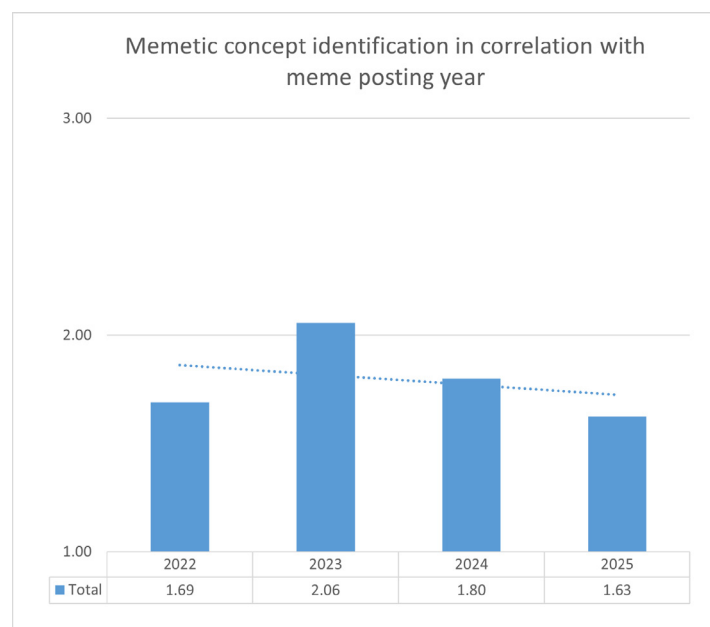


Figure 13: Memetic concept identification in correlation with the meme posting year.

Event or news item referenced

In this test, ChatGPT was asked to identify an event mentioned explicitly or referred to implicitly in a meme. Most of the memes in the Russo-Ukrainian War meme pool reference an event or a news item of a different scale (from a lady hitting a drone with a jar of pickles from her balcony to troop movements), various levels of specificity (from memes about Kyrylo Budanov's facial expressions or the battle for Avdiivka, to memes discussing Ukrainians' overall resilience in the face of the war) and different referencing approaches (from very direct ones, like memes mocking another sunk Russian warship, to memes indirectly pointing to the overall situation, for example, implications of infrastructure damage). What unites all these referenced events is their currency at the time of the memes' creation.

ChatGPT was able to identify event references in 60.77% of all memes correctly. In 17.69% of the memes, the responses were evaluated as partially correct (assigned a score of 2), indicating that they lacked specificity or had a somewhat shifted focus. 21.54% of responses were wrong.

The overall event identification accuracy was lower in the higher, context-heavy end of the CS spectrum (**Figure 14**). In memes built on a non-contextual image or a combination of contextual and non-contextual images, the presence of a memetic idea appears to have a negative effect on the score. In contrast, for the contextual image memes, that effect was positive.

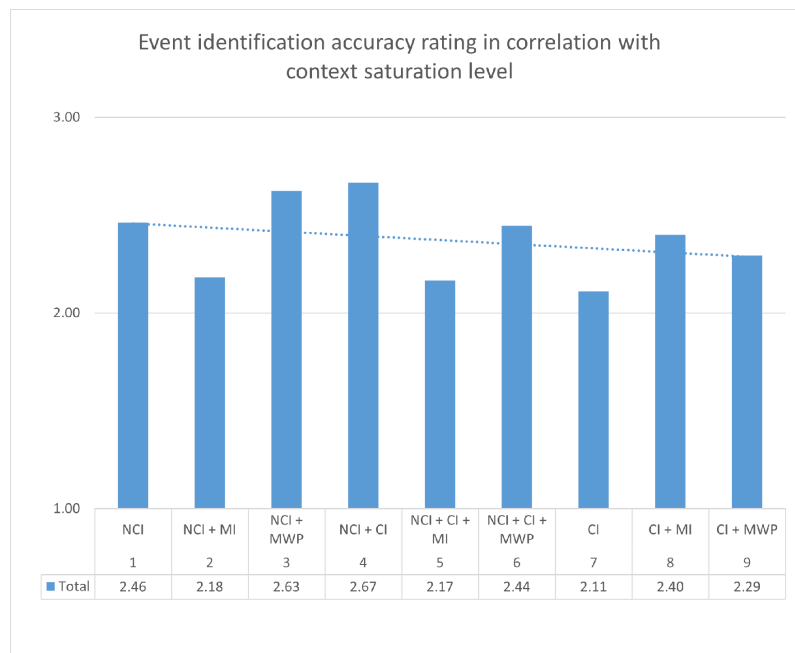


Figure 14: Event identification accuracy in correlation with context saturation level.

ChatGPT showed relatively even performance in event identification accuracy from 2022 to 2024, with a decline for 2025 memes (**Figure 15**).

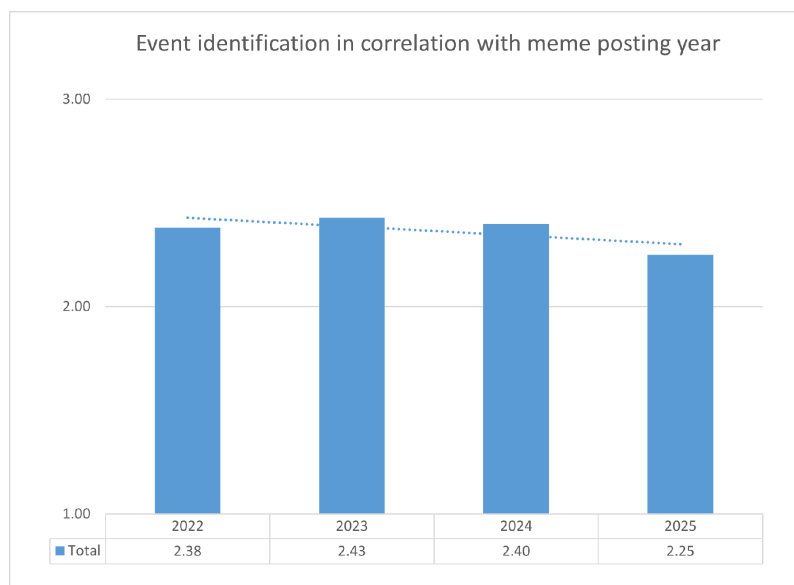


Figure 15: Event identification in correlation with the meme posting year.

Real person featured or mentioned

We asked ChatGPT to identify a person either explicitly mentioned or implied in each meme. This testing category produced the highest performance score (86.15% responses were assessed as correct, 6.15% partially correct, and 7.69% as incorrect).

ChatGPT demonstrated the highest levels of performance in identifying people in memes built on non-contextual images supported by a memetic idea and those built on a combination of contextual and non-contextual images in combination with a memetic idea or a memetic word/phrase (the two tripartite templates, **Figure 16**). The overall performance declines towards the higher context segment of the CS spectrum.

Most of the mistakes by ChatGPT in this test were made in situations where a person was implied indirectly, for example:

- An image of a scooter in memes stands for the assassination of Igor Kirillov, a Russian Army officer killed by an exploding scooter;
- A picture of a long table or a tsar complaining about the need to use a suitcase as a toilet both reference Putin and his alleged outrageous security measures;
- Moses promising his people '2–3 weeks' of wandering makes fun of Oleksiy Arestovych's rhetoric.

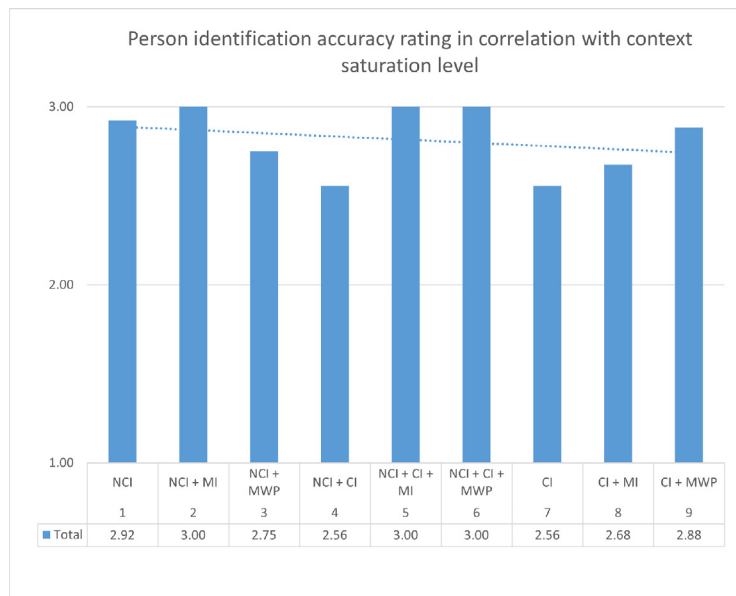


Figure 16: Person identification accuracy in correlation with context saturation level.

In some memes, ChatGPT did not recognize people either because of the untidy mashup work (e.g. V. Zelensky's head photoshopped onto the Disaster Girl, was described as 'likely a photoshopped face of a Ukrainian influencer'), or simply because of the lack of information (e.g. the chatbot failed to recognize the Russian pop-star propagandist Yaroslav Dronov (Shaman) in a popular publicity shot).

Diachronically, ChatGPT demonstrated higher accuracy in identifying people in memes posted in 2023 and especially 2024 (**Figure 17**). The accuracy dropped for the newest 2025 memes.

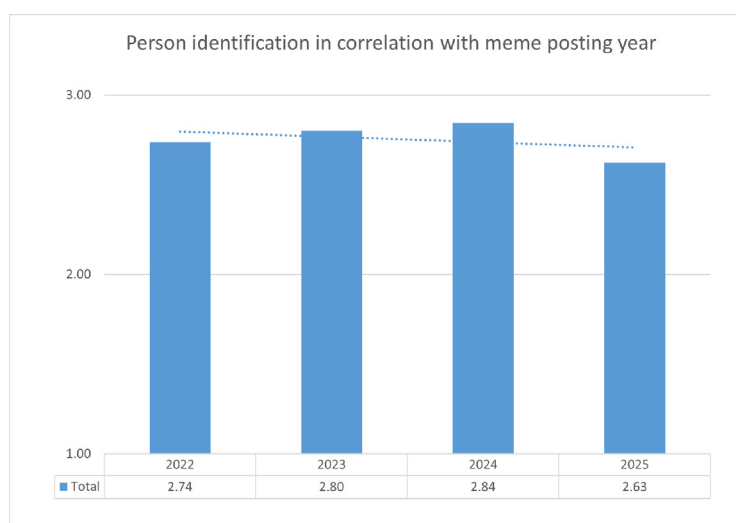


Figure 17: Person identification in correlation with the meme posting year.

Conclusions

We will now summarize ChatGPT's overall performance over the course of our study and draw some conclusions.

In the course of our study, we asked ChatGPT to identify memes' structural elements and write a narrative description of the comic effect. An average score of all seven testing categories was calculated to obtain a combined rating of the chatbot's performance analyzing each meme (**Figure 18**). Overall, we observe a decline in the rating towards the higher CS end of the spectrum, which confirms our general hypothesis that memes with a higher context concentration would present more difficulty for analysis. We also see the most balanced performance in that segment. It appears that cumulative performance follows a specific pattern in all three segments of the CS spectrum: the lowest score was assigned to analysis of memes built only on a visual template, with no MI or MWP. It appears that the presence of an MI and an MWP improves Chat's performance even if it struggles with contextual image template recognition.

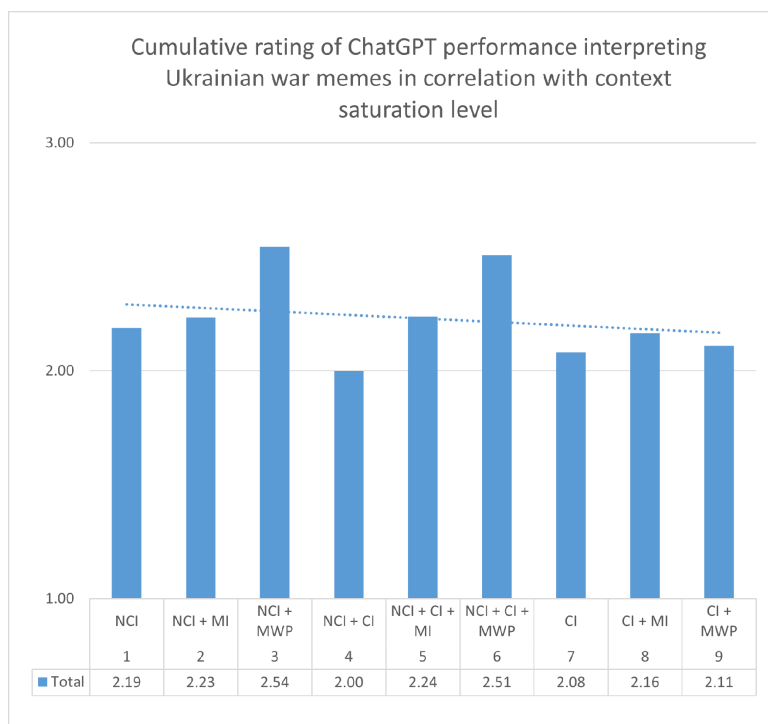


Figure 18: Cumulative score of ChatGPT performance interpreting Ukrainian war memes in correlation with context saturation level.

The CS patterns that ChatGPT handled with the most accuracy were NCI + MWP (level 3 on our CS scale) and NCI + CI + MWP (level 6). The lowest performance was recorded for the memes based on a combination of NCI + CI (level 4) and CI alone (level 7).

When average performance scores in each category are compared, the data shows that ChatGPT performed best in identifying the two tangible context-based elements of the meme: persons featured (average score 2.78 out of 3) and visual templates used (2.62, **Figure 19**). Both of these elements are often presented explicitly and require primarily visual processing. More abstract elements were identified with greater difficulty, ranging from events or news items referred to in a meme (2.32) to viral words/phrases used in a meme (1.74). The verbal templates emerged as the most challenging element for ChatGPT to analyze. In the next stage of this research, it would be reasonable to attempt to improve the Chatbot's performance by presenting it with sets of verbal and cultural references. The present study was designed to test the Chatbot's raw ability to identify them.

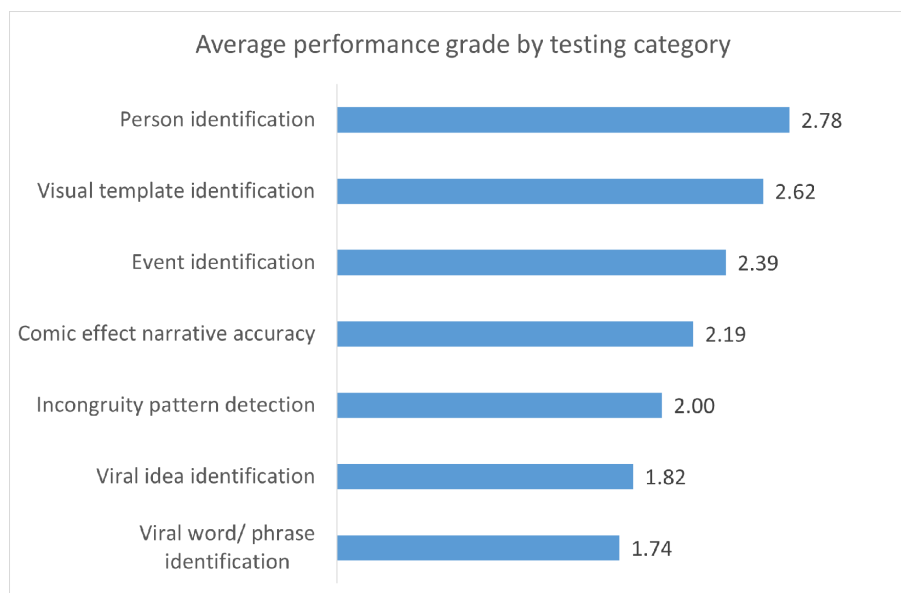


Figure 19: Average performance score by testing category.

We also compared the overall performance rating across the four incongruity patterns: frame, discourse, image-frame, or discourse-image based (coded, respectively, FB, DB, FIB, and DIB). The results (**Figure 20**) show that frame-based memes are the easiest for Chat to interpret (average score 2.39), and discourse-based ones are the most difficult (2.05). It appears that adding an image as an incongruity-forming element improves the bot's understanding of both the discourse-based and the frame-based memes' humor. This is consistent with the Chatbot's demonstrated overall difficulty in interpreting verbal references.

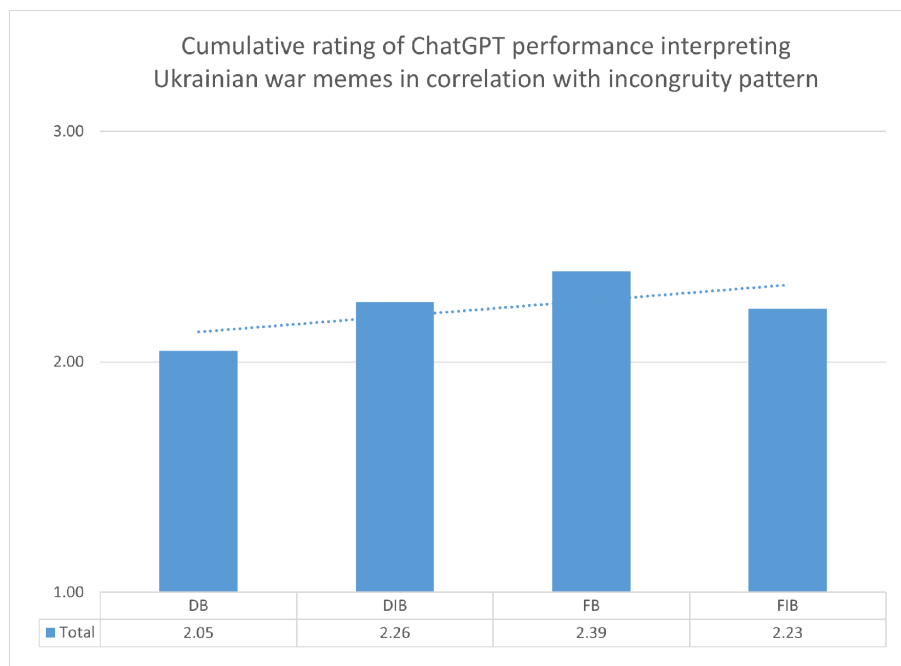


Figure 20: Cumulative rating of ChatGPT performance interpreting Ukrainian war memes in correlation with the incongruity pattern.

We also investigated a possible correlation between ChatGPT's performance in identifying specific contextual elements of the meme structure and content and the accuracy of the comic effect narrative (a prototype for a descriptive metadata element of meme archiving). We discovered that in our sample, this correlation is, in fact, direct for most variables: the more accurately the chatbot identifies specific structural elements, the better the narrative it writes, with two exceptions (**Figure 21**). There is a slight dip in identifying the visual template for the narrative score 2, which then rises again to become the chatbot's second most successful skill (behind person identification). Strikingly, the viral word/phrase identification score falls dramatically from an average of 1.76 for partially correct narratives to 1.64 without affecting entirely correct narratives.

Finally, assessing ChatGPT's ability to analyze retrospective information accurately was a very important area of testing. Positive results could strongly support using AI to generate archival metadata for memes and other born-digital multimodal and multi-textual documents, if it is demonstrated that, unlike the human brain, artificial intelligence retains more rare and obscure information over time. Our test results indicate a generally consistent performance from 2022 to 2025, with minor increases in 2023 and 2024 (**Figure 22**). However, as previously noted, performance on specific structural elements like memetic concept or event identification declined in 2024.

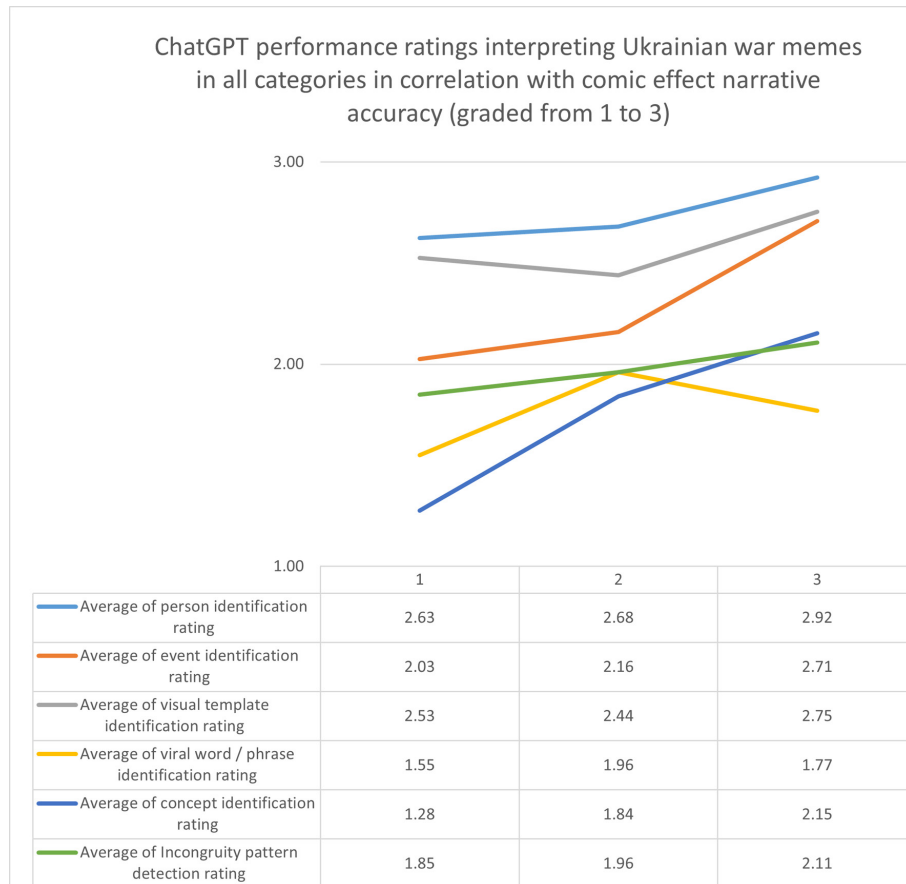


Figure 21: ChatGPT performance interpreting Ukrainian war memes across all categories correlated with comic effect narrative accuracy (ranked from 1 to 3).

Whether these results reflect the development trajectory of ChatGPT itself (notably launched in 2022, coinciding with the Russian invasion and the earliest year of our testing materials: Marr, 2023), showing an uneven performance curve (Chen, Zaharia, and Zou, 2024), the influence of the Ukrainian AI training dataset's development, or the dynamics of the broader online media coverage of the Russo-Ukrainian War in 2023 (Statista, 2025), remains a question for future research.

The results of the study show clearly that there is a direct correlation between ChatGPT's ability to interpret memes and its access to contextual information. In general, the study revealed that interpreting verbal humor (extracting information from memetic words and phrases and analyzing the structure of discourse-based memes) was the most challenging task for ChatGPT, while recognizing persons mentioned either explicitly or implicitly was the easiest. We also discovered that memes with a more complex template structure (that includes not only a memetic image, but also a memetic concept or a verbal meme) are easier for the Chatbot to interpret than simpler memes based on an image template alone.

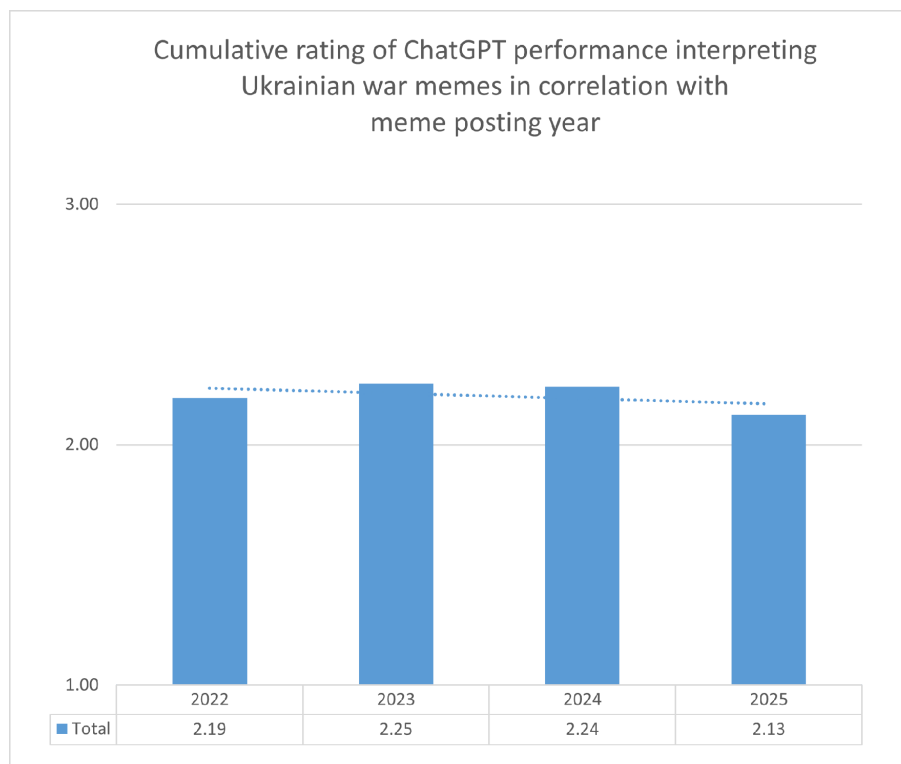


Figure 22: Cumulative score of ChatGPT performance interpreting Ukrainian war memes correlated with meme posting year.

We strongly believe that internet memes must be archived thoughtfully and with meaningful metadata that describes both explicit and implicit elements. Currently, generative AI is not capable of creating detailed metadata for two-dimensional internet memes without human help. It is important to remember that accuracy standards for archival metadata are very high, and although we used a three-level assessment system, only the correct results (score 3) are practically usable. All other results will distort a meme's meaning and effectively end the meme's life.

These conclusions should not be discouraging. We firmly believe that generative AI will be able to play a key role in archival preservation of digital ephemera in the future. It is important to remember that AI technology is evolving quickly and therefore, the results of this study will ultimately become only a snapshot of a specific moment in AI's history.

Competing Interests

The author has no competing interests to declare.

References

- Alam, F, Biswas, Md. R, Shah, U, Zaghoulani, W and Mikros, G** 2024 Propaganda to Hate: A Multimodal Analysis of Arabic Memes with Multi-Agent LLMs.' In: Barhamgi, M, Wang, H, Wang, X (eds), *Web Information Systems Engineering – WISE 2024. Lecture Notes in Computer Science*, vol 15440. Singapore: Springer. pp. 380–390. https://doi.org/10.1007/978-981-96-0576-7_28.
- Anderau, G and Barbarrusa, D** 2024 The Function of Memes in Political Discourse. *Topoi* 43 (5): 1529–46. <https://doi.org/10.1007/s11245-024-10112-0>.
- Bellamkonda, S, Lohakare, M and Patel, S** 2022 A Dataset for Detecting Humor in Telugu Social Media Text. In: Chakravarthi BR, Priyadarshini R, Madasamy AK, Krishnamurthy, P, Sherly E, Mahesan S (eds.) *Proceedings of the Second Workshop on Speech and Language Technologies for Dravidian Languages. Association for Computational Linguistics, Dublin, Ireland. Association for Computational Linguistics*. pp. 9–14. <https://doi.org/10.18653/v1/2022.dravidianlangtech-1.2>.
- Boylan, A L** 2020 *Visual Culture*. Cambridge, Massachusetts: The MIT Press.
- Busdriver242** 2022 It Doesn't Look Good For Them. *Reddit*, February 27. https://web.archive.org/web/20250309135438/https://www.reddit.com/r/memes/comments/t2tyvs/it_doesnt_look_good_for_them/ [Last accessed June 20, 2025].
- Cai, Y, Cai, H and Wan, X** 2019 Multi-Modal Sarcasm Detection in Twitter with Hierarchical Fusion Model. In: Korhonen A, Traum D, Màrquez L (eds.) *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics. Association for Computational Linguistics, Florence, Italy. Stroudsburg: Assoc Computational Linguistics-Acl*. pp. 2506–2515. <https://doi.org/10.18653/v1/P19-1239>.
- Chen, K, Feng, A, Aanegola, R, Saha, K, Wong, A, Schwitzky, Z, Ka-Wei Lee, R, O'Hanlon, R, De Choudhury, M, Altice, F L, Khoshnood, K and Kumar, N** 2023 Categorizing Memes About the Ukraine Conflict. In: Dinh, T N and Li, M (eds), *Computational Data and Social Networks*. Switzerland: Springer. pp. 27–38. https://doi.org/10.1007/978-3-031-26303-3_3.
- Chen, L, Zaharia, M and Zou, J** 2024 How Is ChatGPT's Behavior Changing Over Time? *Harvard Data Science Review*, 6(2). <https://doi.org/10.1162/99608f92.5317da47>.
- Chiruzzo, L, Castro, S, Gongora, S, Rosa, A, Meaney, J A and Mihalcea, R** 2021 Overview of HAHA at IberLEF 2021: Detecting, Rating and Analyzing Humor in Spanish. *Procesamiento del Lenguaje Natural*, 67: 257–68. <https://doi.org/10.26342/2021-67-22>.
- Dawkins, R** 1976 *The Selfish Gene*. New York: Oxford University Press.
- Denisova, A** 2019. *Internet Memes and Society: Social, Cultural, and Political Contexts*. New York: Routledge, Taylor & Francis Group.
- Gal, N, Shifman, L and Kampf, Z** 2016 'It Gets Better': Internet Memes and the Construction of Collective Identity. *New Media & Society*, 18 (8): 1698–1714. <https://doi.org/10.1177/1461444814568784>,1710.

- García López, F and Martínez Cardama, S 2020 Strategies for Preserving Memes as Artefacts of Digital Culture. *Journal of Librarianship and Information Science*, 52 (3): 895–904. <https://doi.org/10.1177/0961000619882070.896>.
- Hossain, E, Sharif, O, Hoque, M M, and Preum, S M 2024 Deciphering Hate: Identifying Hateful Memes and their Targets. *arXiv.org*. <http://arxiv.org/abs/2403.10829> [Last accessed June 20, 2025].
- Imgflip* 2025 Chmonya Meme Generator. <https://imgflip.com/memegenerator/484744467/Chmonya> [Last accessed May 17, 2025].
- Karim, Md R, Dey, S K, Islam, T, Shajalal, M, Chakravarthi, B R, M, A K, Durairaj, T, Mandl, T, Murthy, H, and O’Riordan, C 2023 Multimodal Hate Speech Detection from Bengali Memes And Texts. In: *Speech and Language Technologies for Low-Resource Languages: First International Conference, SPELL 2022, Kalavakkam, India, November 23–25, 2022, Proceedings*. Cham: Springer International Publishing, pp 293–308. https://doi.org/10.1007/978-3-031-33231-9_21.
- Kirner-Ludwig, M 2022 Internet Memes as Multilayered Recontextualization Vehicles in Lay-Political Online Discourse. In: Xie, C (ed.), *The Pragmatics of Internet Memes*. The Netherlands: John Benjamins Publishing Company. pp. 145–81.
- Knežević, A 2023 Internet Memes as Heritage in Becoming and Its Problems: The Methodology of Heritization and the Formation of Cultural Memory in Cyberspace. *Glasnik Etnografskog instituta* 71 (2): 275–298. <https://doi.org/10.2298/GEI2302275K>.
- Kononenko, O (ed.) 2024. ‘Minuty Dve Krichal.’ *Ukraintsy Lishili Zhitelia Belgoroda Mashiny, a Potom Eshche i Prevratili Ego v Mem*. *New Voice*, January 5, 2024. https://nv.ua/lifestyle/obstrel-belgoroda-mestnyy-zhitel-poteryal-mashinu-i-prevratilsya-v-mem-50381671.html#goog_rewarded [Last accessed May 17, 2025].
- Kumari, G, Adak, C and Ekbal, A 2024 Mu2STS: A Multitask Multimodal Sarcasm-Humor-Differential Teacher-Student Model for Sarcastic Meme Detection. In: Macdonald, C, Ounis, I, Goharian, N, He, Y, McDonald, G, Tonello, N and Lipani, A (eds.) *Advances in Information Retrieval: 46th European Conference on Information Retrieval, ECIR 2024, Glasgow, UK, March 24–28, 2024, Proceedings, Part II*. 14609. Switzerland: Springer. pp. 19–37. https://doi.org/10.1007/978-3-031-56063-7_2.
- Laineste, L and Voolaid, P 2016. Laughing across Borders: Intertextuality of Internet Memes. *European Journal of Humour Research*, 4 (4): 26–49. <http://doi.org/10.7592/EJHR2016.4.4.laineste>.
- Lynch, M P 2022 Memes, Misinformation, and Political Meaning. *The Southern Journal of Philosophy* 60 (1): 38–56.
- Marr, B 2023 A Short History Of ChatGPT: How We Got To Where We Are Today. *Forbes*, May 19. <https://www.forbes.com/sites/bernardmarr/2023/05/19/a-short-history-of-chatgpt-how-we-got-to-where-we-are-today/> [Last accessed May 7, 2025].
- Meme-arsenal* 2025 Create Comics Meme ‘Chimpanzee, Monkey with Lips, Chimpanzees’. <https://www.meme-arsenal.com/en/create/template/9203797> [Last accessed May 10, 2025].
- Milner, R M 2016 *The World Made Meme Public Conversations and Participatory Media*. Cambridge, Massachusetts: The MIT Press.
- Nahon, K and Hemsley, J 2013 *Going Viral*. Cambridge: Polity.

- Nyzovets', A** 2023 I Ak Feik pro Portret Bandery Zrobyv Populiarnym Telegram-bot 'Papka v Pakietie'. *LIGA.Life*, January 2, 2023. <https://web.archive.org/web/20230103182715/https://life.liga.net/istoriyi/article/kak-feyk-o-portrete-bandery-sdelal-populyarnym-telegram-bot-papka-v-pakete> [Last accessed May 17, 2025].
- Oboz.ua** 2017 'Koly Vychyslyv za IP': Sotsmerezhi Rozsmishylo Foto z Hroismanom. March 31. <https://news.obozrevatel.com/ukr/society/71496-koli-obchisliv-po-ip-sotsmerezhi-rozsmishilo-foto-z-grojsmanom.htm> [Last accessed September 30, 2025].
- Phillips, W** 2015 *This Is Why We Can't Have Nice Things: Mapping the Relationship between Online Trolling and Mainstream Culture*. Cambridge, Massachusetts: The MIT Press.
- Prabhakaran, V, Qadri, R and Hutchinson, B** 2022 Cultural Incongruencies in Artificial Intelligence. *arXiv:2211.13069*. <https://doi.org/10.48550/arxiv.2211.13069>.
- Rakityanskaya, A** 2025, Archiving Internet Memes of the Russian Invasion of Ukraine: Issues in the Preservation of Born-digital Ephemera and the Case of the SUCHO Meme Wall. *Access Points*, University of Münster Research Group 'Access to Cultural Goods in Digital Transformation', [forthcoming, draft available on request].
- Rossi, C and Bondarenko, I** 2024 Internet Memes: A Cognitive Approach to the Issue of Semantic Translatability. *ILCEA*, 2024: 53. <http://journals.openedition.org/ilcea/19860> [Last accessed June 20, 2025]; <https://doi.org/10.4000/ilcea.19860>.
- Sharma, C, Bhageria, D, Scott, W, PYKL, S, Das, A, Chakraborty, T, Pulabaigari, V, and Gamback, B** 2020 Semeval-2020 Task 8: Memotion Analysis—the Visuo-lingual Metaphor! *arXiv*. <https://doi.org/10.48550/arXiv.2008.03781>.
- Shifman, L** 2014 *Memes in Digital Culture*. Cambridge, Massachusetts: The MIT Press.
- Soriano-Gonzalez, L and Belda-Medina J** 2024 Exploring Image-Text Combinations in Visual Humour through Large Language Models (LLMs). *Digital Scholarship in the Humanities*, 40(1): 280–294. <https://doi.org/10.1093/llc/fqae068>.
- Statista** 2025 Russia-Ukraine War Online News Coverage by Keyword, February 3, 2025. <https://www.statista.com/statistics/1344628/russia-ukraine-war-mentions-in-online-press/> [Last accessed May 5, 2025].
- Suls J M** 1972 A Two-Stage Model for the Appreciation of Jokes and Cartoons: An Information-Processing Analysis. In: Goldstein, J H and McGhee, P E (eds.) *The Psychology of Humor; Theoretical Perspectives and Empirical Issues*. New York, Academic Press. pp. 81–100. <https://doi.org/10.1016/b978-0-12-288950-9.50010-9>.
- Syromiatnikov, M, Ruvinskaya, V and Troynina, A** 2024 ZNO-Eval: Benchmarking Reasoning Capabilities of Large Language Models in Ukrainian. *Informatics. Culture. Technology*, 1 (1): 185–191. <https://doi.org/10.15276/ict.01.2024.27>.
- Thammastitkul, A** 2023. Assessing the Effectiveness of Image Recognition Tools in Metadata Identification through Semantic and Label-Based Analysis. *International Journal of Metadata, Semantics and Ontologies* 16 (3): 227–37. <https://doi.org/10.1504/IJMSO.2023.137174>.
- Thapa, S, Shiwakoti, S, Bikram Shah, S, Adhikari S, Veeramani, H, Nasim, M, and Naseem, U** 2025 Large Language Models (LLM) in Computational Social Science: Prospects, Current State, and

Challenges. *Social Network Analysis and Mining* 15 (1): 4. <https://doi.org/10.1007/s13278-025-01428-9>.

The Guardian 2022 Ukrainian Woman Offers Seeds to Russian Soldiers so 'Sunflowers Grow When They Die' – Video, February 25. <https://web.archive.org/web/20230622024606/https://www.theguardian.com/world/video/2022/feb/25/ukrainian-woman-sunflower-seeds-russian-soldiers-video> [Last accessed June 20, 2025].

Vidomenko, D 2022 Shcho Oznachaie "Bavovna": Poiasnennia Viis'Kovoho Internet-Fenomeny i laskravy Memy. *Glavred*, August 20. <https://news.glavred.net/chto-oznachaet-bavovna-obyasnenie-voennogo-internet-fenomena-i-yarkie-memy-10402130.html> [Last accessed October 10, 2025].

Wang, J, Luo, J, Yang, G, Hong, A, Luo, F, Wani, M, Gama, J, Boicu, M, Abreu, P and Sayed-Mouchaweh, M 2023 Is GPT Powerful Enough to Analyze the Emotions of Memes? In: Wani, M A (ed.) *International Conference on Machine Learning and Applications (ICMLA)*, Los Alamitos, CA: IEEE. pp. 1338–1343. <https://doi.org/10.1109/ICMLA58977.2023.00202>.

Wiggins, B E 2019 *The Discursive Power of Memes in Digital Culture: Ideology, Semiotics, and Intertextuality*. New York: Routledge.

Yus, F 2017. Incongruity-Resolution Cases in Jokes. *Lingua*, 197: 103–122. <https://doi.org/10.1016/j.lingua.2017.02.002>.

Yus, F 2021 Incongruity-Resolution Humorous Strategies in Image Macro Memes. *Internet Pragmatics*, 4(1): 131–149. <https://doi.org/10.1075/ip.00058.yus>.

