

Not Just Novelty: A Longitudinal Study on Utility and Customization of an AI Workflow

Tao Long
long@cs.columbia.edu
Columbia University
New York City, New York, USA

Katy Ilonka Gero
katy@g.harvard.edu
Harvard University
Cambridge, Massachusetts, USA

Lydia B. Chilton
chilton@cs.columbia.edu
Columbia University
New York City, New York, USA

ABSTRACT

Generative AI brings novel and impressive abilities to help people in everyday tasks. There are many AI workflows that solve real and complex problems by chaining AI outputs together with human interaction. Although there is an undeniable lure of AI, it is uncertain how useful generative AI workflows are after the novelty wears off. Additionally, workflows built with generative AI have the potential to be easily customized to fit users' individual needs, but do users take advantage of this? We conducted a three-week longitudinal study with 12 users to understand the familiarization and customization of generative AI tools for science communication. Our study revealed that there exists a familiarization phase, during which users were exploring the novel capabilities of the workflow and discovering which aspects they found useful. After this phase, users understood the workflow and were able to anticipate the outputs. Surprisingly after familiarization, the perceived utility of the system was rated higher than before, indicating that the perceived utility of AI is not just a novelty effect. The increase in benefits mainly comes from end-users' ability to customize prompts, and thus potentially appropriate the system to their own needs. This points to a future where generative AI systems can allow us to design for appropriation.

CCS CONCEPTS

• **Human-centered computing** → **Empirical studies in HCI**; **Human computer interaction (HCI)**; User studies; • **Computing methodologies** → *Natural language generation*.

KEYWORDS

workflow, longitudinal user experience, customization, technology appropriation, ownership, familiarization, mental model, AI chains, scaffolding, science communication, novelty, generative AI, LLMs

ACM Reference Format:

Tao Long, Katy Ilonka Gero, and Lydia B. Chilton. 2024. Not Just Novelty: A Longitudinal Study on Utility and Customization of an AI Workflow. In *Designing Interactive Systems Conference (DIS '24)*, July 1–5, 2024, IT University of Copenhagen, Denmark. ACM, New York, NY, USA, 22 pages. <https://doi.org/10.1145/3643834.3661587>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

DIS '24, July 1–5, 2024, IT University of Copenhagen, Denmark

© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 979-8-4007-0583-0/24/07...\$15.00

<https://doi.org/10.1145/3643834.3661587>

1 INTRODUCTION

Generative AI has impressive abilities to generate text, images, and code based on simple prompts. Building on this, a powerful technique to solve more complex problems is to break problems into AI workflows that chain together prompts [105]. Often, the steps of the chain involve interfaces with human interaction to check the results and steer progress in productive directions. This can enable more complex processes such as following divergent and convergent thinking [41, 100], the human-centered design process [58, 72, 99], the process model of writing [25, 65, 112], and simulating thoughts, behaviors and actions of a group of people interacting [71]. Studies of AI workflows tend to show that people find them useful. They save time [61] and mental effort [65, 72], and increase output [58].

However, most of these systems test their utility in a single lab session. The presence of AI in a workflow brings great excitement, expectations, and hype. But could this just be an effect of the novelty of including AI in a workflow? For workflows that provide scaffolding, it is natural to assume that scaffolding might not be needed after the user has internalized the workflow. Workflows are often criticized for pushing users towards similar outputs and constraining diversity and creativity. Moreover, workflows can be brittle [77]—although they may work for one set of well-defined tasks, they might not generalize to more complex examples that users encounter in the wild. Is the success of AI workflows a product of the novelty effect, or is there a lasting benefit?

The novelty effect is a common effect seen in technology adoption [79, 84, 89] where the desire to use a new technology is initially high but then diminishes over time — as users become more familiar with the technology it seems less magical. Previous studies show that some technologies overcome the novelty effect [40, 79], but some do not [84, 89]. Even though there is a novelty effect, the familiarization effect can actually increase engagement over time [79]. Technology that becomes familiar can become internalized and thus used with less effort, leading to greater rewards. Long-term studies on education technologies indicate that technologies that cater to users' needs and feelings of autonomy can overcome the initial novelty effect [40]. Currently, it is unclear whether the appeal of AI workflows stems from a novelty effect or if there is a lasting benefit past the workflow familiarization.

We conducted a three-week longitudinal study with 12 users to understand the familiarization and customization of generative AI tools for science communication. We study a 7-step human-AI workflow to help STEM experts write tweets that motivate a topic in their field to the general public [61]. The science communication task is called a Tweetorial hook [26], which is challenging because it requires STEM experts to think about examples or experiences with the topic the public might relate to. For example, the web

programming topic of AJAX is very technical, but users have all experienced it on social media when new posts load as they scroll without users explicitly requesting them. A previous study [61] shows that in a single session, a workflow based on a large language model (LLM) reduced the mental demand for writing a Tweakorial hook compared to writing without it (simply using web search). We recreated this workflow and added two simple customization features. First, we display the LLM prompt used at every stage and allow users to edit it to customize the responses and the workflow. Second, we add the hooks written by the users to the prompts as training examples so the system can be personalized based on examples. Then, we study a three-week deployment of the system as design probe [37], over which participants use the system in 10 different sessions to write ten different hooks. We identify if and when familiarization occurs. To understand whether there is a novelty effect, we use surveys, observations, and interviews to assess the system's usefulness before and after familiarization.

Our study revealed that:

- A familiarization phase exists and lasts for a few sessions (4.27 sessions, during which the system is already perceived as useful, 5.47/7).
- Before the familiarization concluded, users were exploring the novel capabilities of the workflow and discovering which aspects they found useful.
- After familiarization, the system is rated more useful, indicating that the novelty does not wear off (a 12.1% improvement, from 5.47/7 to 6.32/7, p -value < 0.005).
- This change in the system's usefulness can mainly be attributed to the increased usefulness of the prompt editing feature (an 11.4% improvement, from 5.07/7 to 5.87/7, p -value < 0.005). The other features show no change in their usefulness scores throughout the study.
- Users edit prompts for three major reasons: to keep generations relevant, to guide narrative directions and styles, and to make the workflow more efficient. Chaining users' previous hooks to the prompt as training examples shows no significant benefits.
- After familiarization, user's mental models of the workflow did not change, but their perceptions of ownership increased. There were nuanced ways that users experienced ownership after multiple experiences with the system, but there is a clear relationship between involvement and ownership.

We conclude by discussing the potential reasons that the novelty did not wear off and by considering its implications for other AI workflows.

2 RELATED WORK

2.1 Generative AI-Based Workflows

Generative AI models hold huge potential for supporting various tasks with their domain knowledge and contextual understanding capabilities [9, 70]. For example, users can use LLMs for simple tasks like answering math questions and complex ones such as reviewing and editing academic papers or even writing screenplay scripts [9, 65]. However, LLMs have many limitations. First, LLMs

face challenges in producing high-quality outputs when LLMs have limited knowledge about the task [20] or multi-step reasoning is required [101]. Second, due to the limitations of the underlying generative models, LLM outputs often contain hallucination, repetitiveness, and vague language [9, 35, 38]. Third, as LLMs can take any natural language prompts as inputs, users may find it challenging to determine how to engage and collaborate with the LLMs to troubleshoot their prompts for enhanced outputs [105].

AI workflows help solve complex real-life problems and mitigate LLM shortcomings. As the foundation structure of workflows, AI chains [105] (or LLM chaining) break down complex tasks into simpler steps then chain the LLM output of one step to the input of the next LLM step [29, 104, 105]. This multi-step process helps complete complex tasks that LLMs are not able to do well in one prompt. Also, chaining allows users to easily evaluate the outputs within each step to identify errors, thus providing user control to find errors, fix them, and improve generations according to their preferences [29, 61, 106].

AI workflows successfully complement users with guided support in many traditionally cognitively demanding tasks. For example, various workflows demonstrate their strong system performances or satisfactory user experiences to support writing [52, 61, 65, 74, 85, 87, 99, 109], brainstorming [11, 13, 43, 72, 86, 94, 100, 106, 112], visual design [21, 36, 41, 44, 57, 58, 73, 92, 97, 107] and question answering and sensemaking [1, 18, 19, 24, 42, 53, 95, 108, 114].

However, most of these studies are conducted solely within the framework of controlled user studies lasting typically 30 to 60 minutes. This approach might only capture users' initial reactions of satisfaction or dissatisfaction with the system rather than providing more comprehensive experiences [8, 28, 39, 93]. Thus, a gap exists in understanding how AI workflows perform and whether they continue to aid users well after repeated usage.

2.2 Novelty Effect

The novelty effect is a common effect seen in technology adoption [14, 47, 50], where the desire to use new technology is initially high, but then diminishes over time. It is related to humans' innate desire for novelty and our psychological stress response: individuals exhibit a stronger stress response the first time they face a challenging or threatening experience [34, 82]. With the stress, users are more likely to focus, embrace learning through trials and errors with a higher tolerance, and perceive higher performance [48, 69, 81, 82]. However, novelty effects will eventually diminish as users become more familiar with the technology, with the duration varying depending on the technology [47, 50].

After the novelty wears off, some technologies fail to overcome this effect, wherein users find them less useful after becoming familiar with them, ultimately leading to dissatisfaction and abandonment [60, 84, 89, 98]. However, many technologies successfully navigate this challenge, as users continue to find them useful over time [40, 79]. Also, even though the novelty effect plays a role, familiarizing these technologies can increase users' engagement and perceived usefulness over time [79]. Technology that becomes familiar can become internalized and thus used with less effort, leading to greater rewards. Additionally, long-term studies on education technologies indicate that tools that cater to users' needs and

feelings of autonomy can overcome the initial novelty effect [40]. This is aligned with self-determination theory [15] in educational tools and technology: when the needs for autonomy, competence, and relatedness are satisfied, positive outcomes follow [2, 63, 83].

Generative AI workflows also draw concern regarding the strong novelty effect due to their seemingly dazzling performance [67]. There exists a gap in examining if the benefits of AI tools and workflows hold up and what happens after the novelty wears off.

2.3 Technology Customization and Appropriation

Users customize and appropriate the technology over usage and time. Many long-term studies [4, 30, 33, 46, 51, 54, 90, 113] and technology domestication theory [5] mapped user's journey of adopting and customizing these technologies. Beginning with the learning and comprehending the system affordances, especially the essential features, users navigate the learning phase [14, 76]. Once they better understand the technology and the supported task, users engage in customization to shape the technology for better alignment with their own preference [3, 30, 66, 111]. Then, users start to appropriate the technology by modifying it and applying it for new use cases that the original designers had not anticipated [12, 16, 17, 49, 68]. Along with the increased familiarization and insights gained over usages, users gradually develop their own perceptions, mental models, and opinions towards the system [27, 30, 46, 80]. Technologies that allow customization and appropriation can elevate user satisfaction and performance and meet more user needs [16, 21, 45, 75]. Both customization and appropriation offer users a better sense of control and ownership over the system and its generations [16, 64].

A 2007 paper [16] offers guidelines to design systems that support customization and appropriation, outlining three key principles. First, the system should "allow interpretation": enable users to assign their interpretation to system elements instead of fixing them. Second, it needs to "provide visibility": make the system's functionality transparent for users to understand the effects of their actions. Third, it needs to "support not control" and offer "plugability and configuration": design for efficiency and flexibility, allowing the system to automate and users to combine parts of the system according to their needs.

For generative AI workflows, natural language prompts largely simplify the customization process. Given the flexibility to alter the system through word-level changes, users can easily customize workflows by tweaking prompts or incorporating examples from their previous work. It broadens the access to customization for non-experts [56, 110], in contrast to traditional methods requiring developers and designers to engage directly with users for system updates. While this form of customization holds the potential to enhance the utility of AI systems over time, it remains uncertain whether it can overcome the novelty effect and how users will actually engage in customization or appropriation.

3 BACKGROUND ON TWEETORIAL HOOKS

3.1 Science Communication and Tweetorials

Communicating science effectively and engagingly is key for the public to understand and navigate our rapidly changing world. Historically, the realm of science communication was largely confined

within the boundaries of scholarly papers [10]. These papers are intended for a niche audience of fellow scientists and scholars, and are often difficult for lay people to understand because the writing contains jargon, assumes background knowledge and does not motivate the material or make it relatable [102]. However, with the rise of social media, there has been a marked shift in how people access scientific information online: more and more people are starting to learn science in fragmented time on social media [31, 102].

Starting in late 2017, Tweetorials have become a popular science communication medium for STEM experts to reach the general audience [96]. A 2021 study analyzed popular Tweetorials to identify the science communication principles that made them effective [26]. They found that Tweetorials are generally short (threads of 8 - 15 tweets) and use informal, narrative-driven, and personal language.

As shown in Figure 1, the key feature of a Treetorial is the first tweet, which is called the "hook" to motivate the topic and draw readers in. The hook has an intriguing question: "Are your friends more popular than you?" It provides a relatable idea: "If you check out your Twitter friends, you might find that on average they have more friends than you." This introduces the topic of the Friendship Paradox without jargon. And the details make it really personal as something the reader may have encountered in their own life. Then, the rest of the Treetorial is the "body" that explains the concept, where the hook can pique the reader's interest in continuing to read about this real-world experience.

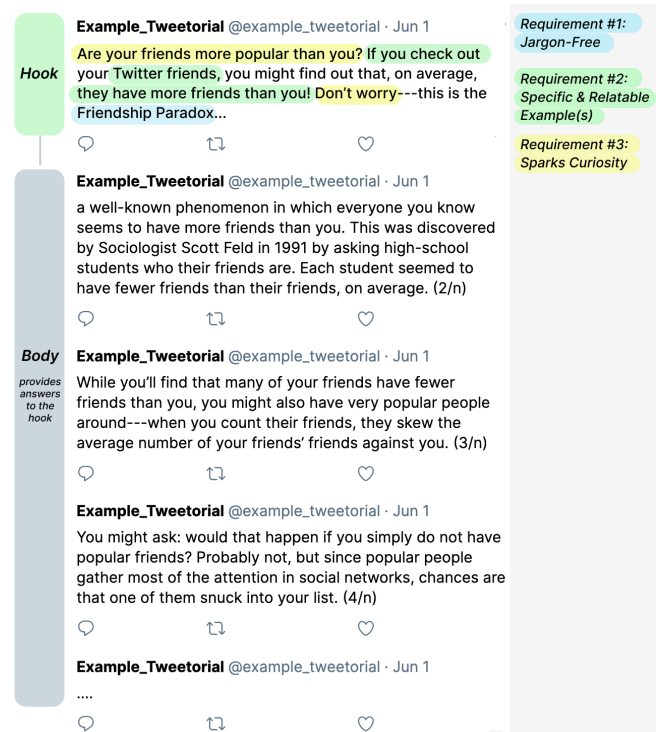


Figure 1: An example Treetorial about friendship paradox, a technical concept in social networks, annotated for its structure and adherence to the requirements (See section 3.2). It comes with a hook and later body sections.

However, studies on Tweetorial writing have shown that writing hooks are a key challenge for STEM experts [26]. They are trained to write about their work in a formal tone for other experts, and it is difficult to go against that training. Another difficulty is finding examples and layman’s terms to motivate the topic for general audiences. Also, they feel uncomfortable using subjective and informal language and avoid personal details, even though 80% of the Tweetorials have them. In an exploratory study using LLMs to support Tweetorial writers, one of the major use cases was ideating concrete examples for the hook [25]. This indicates there is potential to use LLMs to help write Tweetorial hooks.

3.2 What Makes a Good Tweetorial Hook?

Previous work [26] has analyzed Tweetorial hooks using science communication principles, and they found that most hooks had three important characteristics, which were subsequently formulated into requirements for assessing an effective hook by [61]:

- **Requirement 1: Jargon-Free.** The hook should contain no jargon so the general audience can understand it easily, with the possible exception of naming the topic to intrigue readers.
- **Requirement 2: Specific and Relatable Example(s).** The hook should include an analogy, a current event, a popular misconception, or a common experience.
- **Requirement 3: Sparks Curiosity.** The hook should include an intriguing question that sparks readers’ curiosity and drives readers to continue reading.

Figure 2 is an example of Tweetorial hooks for Language Models, a computer science topic, exhibiting these properties:

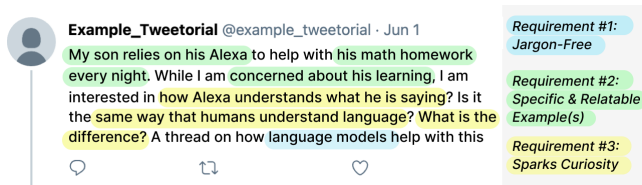


Figure 2: An example Tweetorial hook about language models, a technical concept in computer science, annotated for its adherence to the requirements. In this hook, the relatable experience is talking to a smart device at home, like Amazon Alexa. The specific details include parents’ concerns about the kids’ learning and Alexa helping with math homework. The concerns about his son spark curiosity: “how Alexa understand what he is saying?”

As shown, the relatable and specific content is an experience with the technology. Although this doesn’t work for all topics, it tends to work well for computer science topics because those concepts are usually the technologies we all rely on and can be comparatively smoothly contextualized into real-life use cases.

3.3 Existing AI Workflow for Tweetorial Hooks

A 2023 study [61] presents a human-AI workflow that scaffolds the process of writing a Tweetorial hook. As shown in Figure 3, the

workflow has seven steps and uses AI chains to help with the generations among them. The user first inputs their topic (ex. “*language models*”), and the system generates five everyday examples, and the user selects or edits one (ex. “*autocomplete feature in gmail*”). Based on the topic and example, the system generates five common experiences, and the user picks one (ex. “*helps in finding the next words and writing emails faster*”). Next, the system generates five user scenarios where people might have this experience (ex. “*rushing to email the professor about a homework extension*”). The user can again select or edit one. Next, the system generates a personal anecdote (ex. “*I was sending a last-minute email to my professor asking for extended deadlines...*”). Then, based on it, it continued to contextualize and generate a more specific anecdote about the user scenario, which includes more details to make users feel relatable and personal (ex. “*I was sick... so I’m asking for an extension on submitting my English Literature thesis...*”). Lastly, the system consolidates details from all the previous steps to generate example hooks, allowing users to either use them directly or start from them in the final step to write a hook.

This workflow was then validated by a user study (n=10) comparing the hook-writing experience with and without the system [61]. The workflow was shown to significantly reduce mental demand, effort, and frustration while improving performance. However, this study was conducted in a single 60-minute session, while the novelty of the workflow might have influenced the users’ perception. Therefore, we conducted a study over three weeks to understand users’ longitudinal interaction with the workflow.

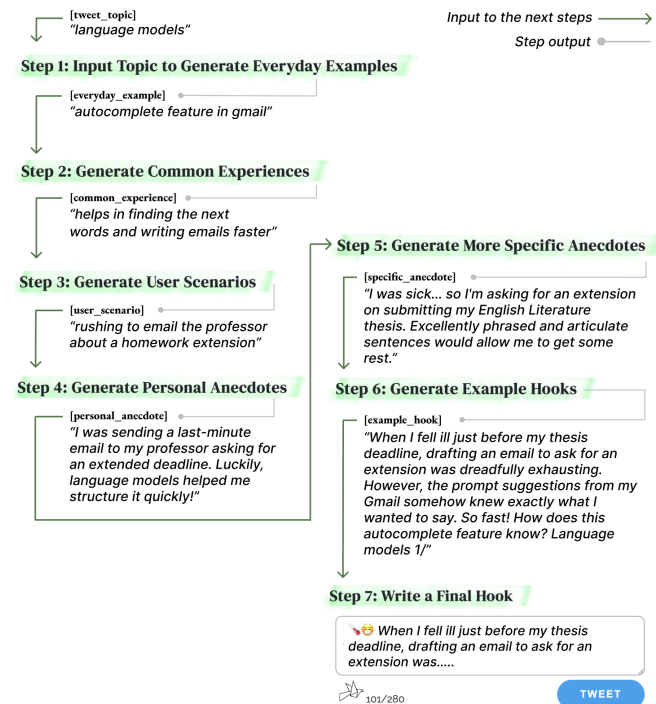
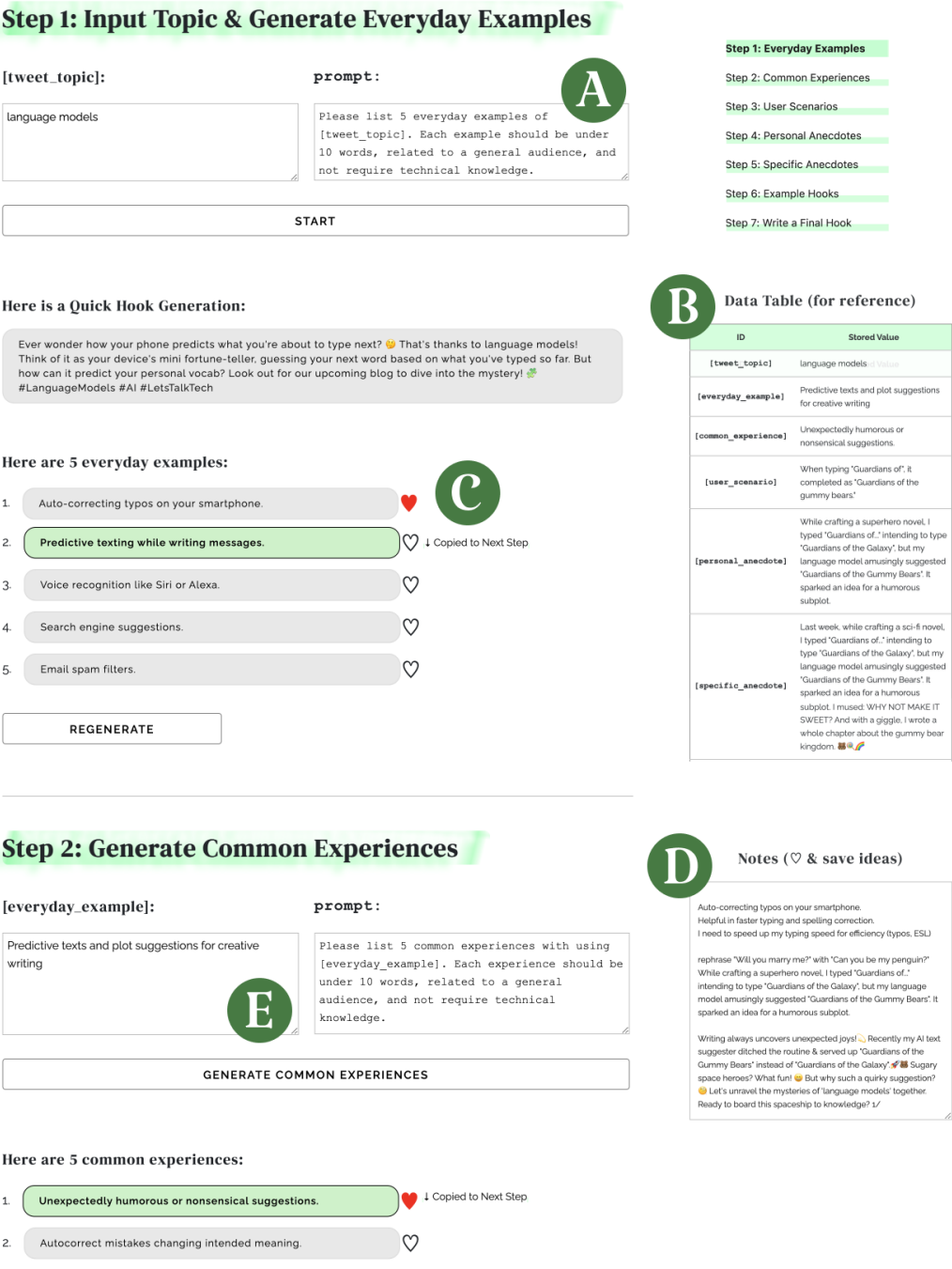


Figure 3: An illustration of the existing AI workflow for writing Tweetorial hooks, following [61]. For each step inside the workflow, users can regenerate, modify, or accept the workflow suggestions before going to the next steps.

4 SYSTEM



4.1 System Design

As shown in Figure 4, we followed the existing workflow in section 3.3 and designed a web application as a design probe [37] to understand users' long-term experience with AI workflows. Specifically, we want to understand (i) whether users want to continue using the full system after repeated usage - whether the novelty wears off, and (ii) whether users will take advantage of the potential of customization in AI workflows.

To address (i), we introduced "Quick Hook Generation" so users see an example hook based on their input topic in Step 1. If users find that the backend model can already fulfill their needs or if the novelty effects of the workflow wear off, users can proceed directly without going through the full system. For (ii), to facilitate user customization throughout usages, we display the system's prompts with variable placeholders, such as [tweet_topic] or [everyday_example], thus empowering users to understand, edit, and customize their system, and promoting user control and autonomy. Starting in the third week (the eighth session), users can choose to feed their previously written hooks back into the system as training examples. We believe these minor adjustments retain the benefits of the original workflow while augmenting the user experience for longitudinal interactions.

We developed an interactive web application prototype utilizing Python, JavaScript, Flask, and OpenAI's GPT-4 API (version-date: Aug-14-2023 to Sep-3-2023). We logged all user interactions and system generations.

5 STUDY

To better understand novelty effects, familiarization and customization in generative AI workflow, this study investigates the following research questions (RQs):

- **RQ 1: Familiarization.** How long is the familiarization phase? How do the usefulness and performance of the system change before and after familiarization?
- **RQ 2: Prompt Editing.** How and why do users customize prompts?

- **RQ 3: Mental Model and Ownership.** How do users' mental models, sense of ownership, and involvement with the system evolve over time?

5.1 Participants

We recruited PhD-level researchers in computer science (CS) who were interested in communicating their research to the public on social media. Because we wanted domain experts, we targeted senior PhD students, but anyone with at least two years of research expertise was deemed expert enough. We focused on CS as a discipline because the original workflow was shown to work for CS topics. Other fields like microbiology or chemistry were not as easy to motivate with an everyday experience. Other workflows could be designed for those fields, but to test the long-term effects of a workflow, we focused on a single domain where it has already been shown to work.

Recruitment. We advertised to PhD students in a CS department through department-wide emails and Slack communications. The advert described that the study would last three weeks, would require about six hours of time in total, and that participants would be compensated each week - for a total of \$150 if they completed all three weeks. All recruitment materials and processes were approved by the university IRB.

Demographics. As listed in Table 1, we had 12 PhD-level students in CS to complete the three-week study. Originally, we recruited 13 participants, but P5 dropped out after day one due to personal circumstances. Out of the final 12 participants, the average age was 27.5, with six men and six women. On average, they had five years of research experience. Their expertise spans various CS research areas, including theory, systems, security, machine learning, natural language processing, and human-computer interaction. They have varying levels of experience with large language models: six out of 12 users use it weekly, four use it monthly, one uses it daily, and one has never used it. All users had frequent science communication experience, mostly at conferences. Only five participants mentioned that they tweeted about scientific topics or papers.

Table 1: Overview of participant demographics and research background

Participant ID	Gender	Age	Research Area	PhD Year	Research Experience	SciComm Frequency
P1	Male	39	Security and Reliability of Systems	5th Year	5 Years	Weekly
P2	Male	25	Theoretical Computer Science	4th Year	6 Years	Weekly
P3	Male	26	Software Systems	2nd Year	4 Years	Monthly
P4	Female	24	Natural Language Processing	1st Year	4 Years	Weekly
P5 (Dropped)	-	-	-	-	-	-
P6	Female	28	Speech Processing	2nd Year	5 Years	Weekly
P7	Male	26	System for Machine Learning	3rd Year	5 Years	Daily
P8	Female	24	Natural Language Processing	2nd Year	2 Years	Daily
P9	Male	23	Natural Language Processing	1st Year	5 Years	Weekly
P10	Female	26	Natural Language Processing	4th Year	7 Years	Daily
P11	Female	34	Human-Computer Interaction	N/A	3 Years	Weekly
P12	Male	24	Virtual Reality	1st Year	2 Years	Monthly
P13	Female	31	Machine Learning Theory	4th Year	10 Years	Daily

5.2 Procedure

As shown in Figure 5, all 12 participants wrote 10 hooks in 10 separate sessions over the course of three weeks. Each session lasted around 20 minutes, followed by a 10-minute survey. In the first session, users are introduced to the system and study and write a hook with the system with the study administrator present to observe the process and answer questions about the task, the system, and the study. In the remaining nine sessions, the following pattern is repeated three times: two solo system usages (without the administrator present), then a system usage with the administrator present to observe and conduct a 30-minute semi-structured interview.

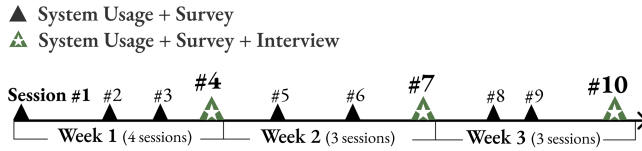


Figure 5: Timeline of study activities over three weeks. Each week, users need to use the system three times to write Tweeporial hooks, followed by a semi-structured interview conducted at the week's end.

5.3 Data Collection and Analysis

System Log. We logged all 120 system usage sessions to record all interaction details, including the user's inputs and prompts for each step, the system's generation results, the system outputs they bookmarked and proceeded with, and the timestamps associated with each activity. During each system usage (either personal use or with the study administrator), we logged every interaction on the Python server backend in a separate .log file. After the user concluded the usage session, a JSON file with the same information was automatically downloaded to the user's laptop, allowing users to review their experiences if needed and ensuring transparency in data saving. Then, they need to submit the file for the session survey to verify their usage.

To analyze the log data, the first author used Python to analyze the user's session interactions, including extracting the duration spent on the system and each step, a list of steps taken, skipped, or revisited by users, and the input data and prompts utilized for each step's generation. Then, to understand how users make prompt edits, the first author used a collaborative affinity diagramming tool to group users' new prompts into subgroups and groups based on their editing objectives with an inductive approach. The final (sub)groups and their examples were iteratively refined and discussed, with feedback sought from the third author.

Survey. We received 118 surveys for the 120 system usage sessions¹ to understand users' task experience, perceptions of the system's usefulness, and performance. After each system usage session, users were asked to complete a Google Form survey. After submitting their participant ID and their session JSON files, users

¹Each user was supposed to submit 10 surveys, resulting in a total of 120 surveys. However, P10 and P13 missed their survey submission for their first usage sessions. The study administrator later followed up with them about their experiences during week one interviews, and both users shared they didn't notice any significant difference between their first and second usage.

needed to fill out three sections: Section A focused on the hook writing task experience alone and included questions adapted from the NASA Task Load Index (NASA TLX) [32]. Section B primarily focused on their system experience, probing participants' perceptions regarding the system's usefulness, mental model, and expectations for future uses. Lastly, Section C had users self-assess the quality and effectiveness of their hook, followed up with a question on the level of ownership they felt over it.

To analyze the survey data, the first author conducted exploratory data analysis using Python to see the relationship between the variables and uncover relationships among session experiences, system performance, user perceptions, and end output evaluations over time and novelty. Then, the first author conducted statistical testing to test statistical significance.

Interview. We conducted 36 interviews with users to understand their degree of familiarity with the system and other experiences and thoughts. Before each interview, users first engaged in a system usage session with the study administrator, during which they were encouraged to think aloud, express their feelings, explain their rationales behind specific choices, and detail their overall experience with the system. Following this, we conducted semi-structured interviews to understand their experiences from the past week, covering any noticeable shifts in their experiences, as well as their expectations for the upcoming week regarding workflow and experiences. We audio-recorded and transcribed the interview sessions using Zoom's auto-transcription feature after receiving the user's consent each time.

To analyze the interview data, the first author performed qualitative coding on the interview transcripts, which were then verified by the second and third authors. The qualitative coding follows the inductive thematic approach [22] to discern patterns and shifts in user experiences, building an affinity diagram based on the codes. The first author then iterated on the codes and (sub)themes to answer research questions, reporting to the other two authors more than five times for refinement.

6 RESULTS

Initial User Experience. Based on the initial survey responses and insights gathered from week one interviews, all 12 users had positive experiences with the system during their first system usage session. Specifically, the users already found the Tweeporial hook writing task challenging and found the system useful in supporting the task. Notably, the ratings for the system were already high: the average rating for the writing task experience was 5.72 out of 7, while the perceived system usefulness scored 5.82 out of 7. Nine users also mentioned the system helped them better express themselves during the first week interview. These comparatively high starting ratings aligned well with findings from earlier research on workflow and the novelty effect, as users might find such automated systems new and "magical" only in the early phase. Thus, the following section addresses the following questions: Were these positive reactions primarily driven by a novelty effect? And more importantly, would users' perceptions of the system's usefulness persist over time?

RQ 1: FAMILIARIZATION

How long is the familiarization phase (RQ 1.1)? How do the usefulness and performance of the system change before and after familiarization (RQ 1.2 and RQ 1.3)?

RQ 1.1 How long is the familiarization phase?

Summary: On average, the familiarization ends after 4.27 sessions. Before familiarization, users were exploring to gain a better understanding of the system. After familiarization, users understood the workflow, were able to anticipate the workflow outputs, and finalized their own versions of the workflow. These changes influenced users' perceptions of the system's usefulness and performance.

To understand when users became familiar with the system, we asked them to reflect on their experience and the degree to which they are familiar with the system during our week two and week three interviews (after sessions 7 and 10). During the two interviews, we first asked, "Do you feel you fully understand the workflow now?" Then, if the user shared that they had already mastered the system, we followed up with the question, "At which session do you think you fully understood the systems?" On average, users identified their familiarization phase ending at 4.27 sessions (median: 4 sessions). This equates to 1.42 hours of system usage, and it was achieved in the first half of week 2 of the study (Table 2 shows the specific session number for each user needed to end familiarization).

We asked the users to characterize their experiences and interactions before familiarization. All 12 users shared that they were exploring and attempting to understand the system. Four users (P1, P3, P7, and P11) described that they were "going with the workflow" rather than trying to edit or critique it. Similarly, three users (P4, P9, and P13) chose from the options provided with little editing. P8 expressed it as follows:

"I feel like [the familiarization] ends early this week, [so after] the fourth usage. I just remembered what each step is and... understood it in the way, not from the definition of the step, but also from the generation results of the step. So that kind of helped me to climb through the learning curve..."

— P8, Week 2 Interview

After familiarization, users expressed three key changes in their experiences. First, users mentioned they understood each step more concretely and how the steps build towards a final hook. Second, users were able to anticipate system generations. They could more quickly decide which steps to skip and which steps to spend more time on based on their importance to the outcome and the abilities of the LLM. Third, ten users felt they "finalized" their own version of the workflow. They identified their preferred ways of using the system and the prompt edits and strategies that work best for their respective topics. P3 expressed it as follows:

"[Now I] have my own fixed routine to generate hook... For me, the workflow starts to converge... the output is satisfying for me, so I kinda stop exploring a new way to improve the result. I have a fixed workflow... so my workflow stops changing."

— P3, Week 3 Interview

One of the 12 users, P2, did not reach familiarization within the ten sessions of the study. When reflecting on why the system didn't work, they said that their topics were in an area of theoretical computer science, which lacked obvious everyday use cases that the system was designed to support. Then, they thought analogies were a more natural fit for motivating his theoretical topics. Thus, they tried adapting the workflows to generate analogies instead of relying on the use case mechanism and thought that was promising. We return to this challenge in the Discussion section 7.2.

Table 2: Number of sessions until the familiarization phase concludes.

Participant ID	Familiarization concluded by # sessions
P1	5
P2	Didn't conclude familiarization
P3	3
P4	6
P6	4
P7	4
P8	4
P9	4
P10	5
P11	4
P12	5
P13	3
Average	4.27
Median	4

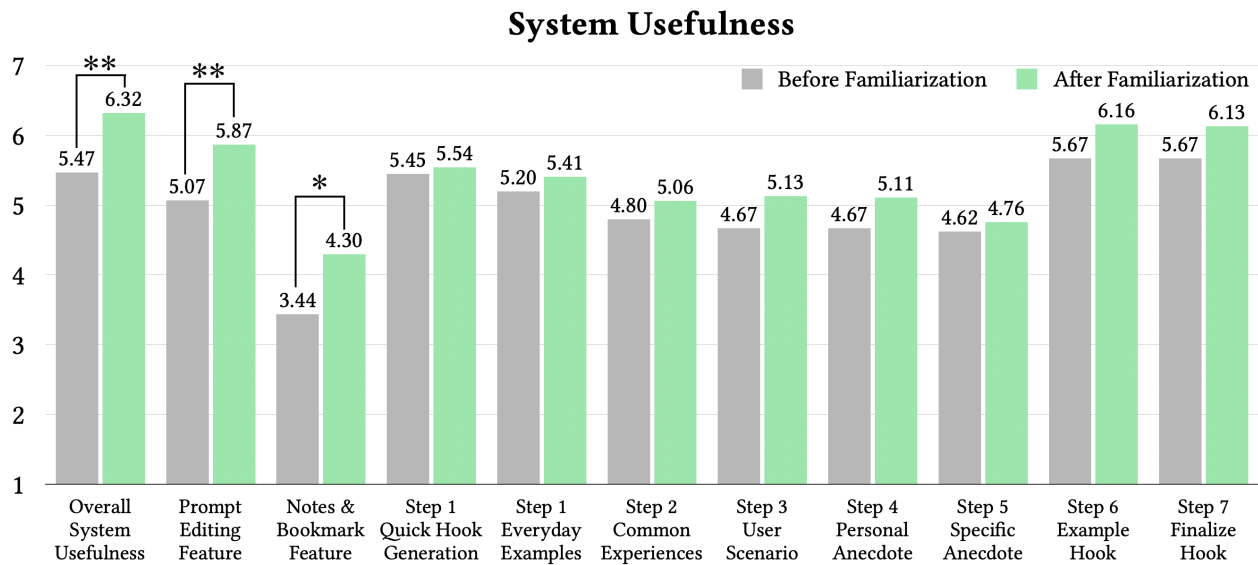


Figure 6: System Usefulness (average scores across 12 users before and after familiarization, ** denotes statistical significance at the p -value < 0.005 level, * denotes statistical significance at the p -value < 0.05 level)

RQ 1.2 System usefulness before and after familiarization²

Summary: On average, the overall system usefulness improved by 12.1% after familiarization. The only features that show an increase in usefulness are the prompt editing feature and the notes & bookmark feature, which serves as a support for prompt editing. In contrast, the usefulness of all the individual steps within the workflow (Step 1 to 7) remained unchanged.

Based on the session survey results shown in Figure 6, users reported a 12.1% (0.85/7) improvement in overall system usefulness before and after familiarization (5.47/7 to 6.32/7, p -value < 0.005). This indicates that there is no novelty effect – over time, workflow utility increases rather than decreases.

When analyzing where the benefits come from, we analyzed the survey on the usefulness of specific system features. On average, there is no change in the perceived usefulness of the individual system steps (Step 1, Step 2,..., Step 7). By week 3, most users still felt every step was “useful and meaningful (P1),” mostly by providing an easy “direction to start (P6)”. At the conclusion of the study, P12 continues to find all the steps and the system useful:

“Even in Steps 1, 2, 3, I’m finding that I need to regenerate a lot to get what I want. It’s still much easier than coming up with everything on my own. [...] I don’t think there’s a single step that I don’t

find useful. Even Step 7 is useful because it gives me like a word counter for tweets.”

— P12, Week 3 Interview

Instead, the system usefulness data showed only two features with statistically higher perceived usefulness: prompt editing with an 11.4% improvement (5.07/7 to 5.87/7, p -value < 0.005) and notes & bookmarks with a 12.3% improvement (3.44/7 to 4.3/7, p -value < 0.05). Both of these features are essential for customizing the workflow. Through prompt editing, users can tailor the system to their style and topic. The bookmarks feature allows people to “heart” their favorite generation, enabling them to explore multiple generations before choosing the best one. As users generate customized prompts and generate more, bookmarks become increasingly important. Thus, it seems that two customization-related features are mostly responsible for the increase in overall system usefulness. For example, P13 noted this association in their week three interview:

“I now better know how to deal with the prompt [editing] and the [bookmark & save feature] box on the left side... It was useful... After learning how I can edit the prompt, yes, [the system gets] more useful.”

— P13, Week 3 Interview

RQ 1.3 Performance before and after familiarization²

Summary: Users reported a 14.9% improvement in their task performance and a 7.1% improvement in time efficiency. From their self-assessments, users believe they produced better hooks, with more specific and relatable examples (Requirement 2) and better at sparking curiosity (Requirement 3).

²Based on each user’s familiarization phase cutoffs, we categorized 120 user sessions into two phases: ‘Before Familiarization’ with 57 sessions and ‘After Familiarization’ with 63 sessions for our log analysis and survey results. We put all ten sessions of P2 under the ‘Before Familiarization’ tag, as P2 indicated that they did not reach familiarization within the entire study.

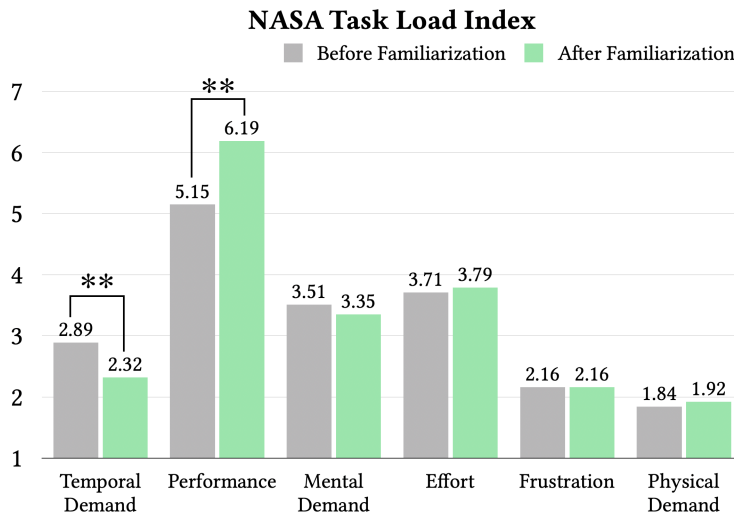


Figure 7: NASA Task Load Index (average scores across 12 users before and after familiarization, ** denotes statistical significance at the p-value < 0.005 level)

From the NASA TLX scores reported from the surveys, users noticed a 14.9% improvement in performance (5.15/7 to 6.19, p-value < 0.005) and a 7.1% improvement in time efficiency for completing the task (2.89/7 to 2.32/7, p-value < 0.005) after familiarization. However, as shown in Figure 7, the task remained similarly mentally demanding, effort-demanding, and frustration-inducing over the two stages, with changes of less than 0.2. Paired t-tests further confirmed that these marginal differences were not statistically significant. This is consistent with RQ 1.1 that people explore less after familiarization, thus saving them time but not necessarily further reducing mental load. For example, P6 shared in the week three that they noticed an improvement in temporal demand:

"I used the system much faster this week... For the first week, I tried to understand what [each step] means... for the last week, I already knew what it means... so I chose the [generation result] option quicker than at the beginning."

— P6, Week 3 Interview

Based on users' self-assessments of their final hooks, they first reported producing better hooks after the familiarization phase. In Figure 8, users shared a 9% improvement for *Requirement 2: Specific and Relatable Example(s)* (from 5.80/7 to 6.43/7, p-value < 0.005) and an 8% improvement for *Requirement 3: Sparks Curiosity* (from 5.69/7 to 6.25/7, p-value < 0.005). Also, users noted a 17.3% reduction in the frequency of hooks containing factual inaccuracies, biases, or potential risks. Before familiarization, 68.4% of hooks were labeled as accurate and unbiased (39 out of the 57 hooks). After familiarization, this increased to 85.7% (54 out of 63 hooks).

This again argues that familiarization with the AI workflow was not just a novelty effect but provided genuine value to users.

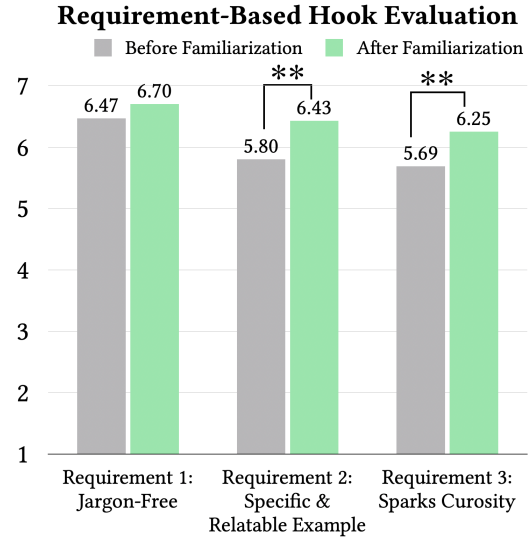


Figure 8: Hook evaluation based on three requirements (average scores across 12 users before and after familiarization, ** denotes statistical significance at the p-value < 0.005 level)

RQ 2: PROMPT EDITING

How and why do users customize prompts?

Summary: Users customized 12.4% of the available prompts (89 of 720). These edits addressed three main problems with the AI workflow: going off-topic, undesired narrative and style, and inefficient workflow. In week 3, we added users' own hooks into the prompt as training examples, but there was no significant increase in the utility of the system. Nine users favored the new feature, while three reported no benefits.

Prompt Editing Themes & Strategies

Across 120 sessions, all users made a collective total of 89 prompt edits. One author coded the 89 edited prompts using the inductive thematic analysis and reported to another author for feedback and iterations on codes. Table 3 shows the six prompt editing strategies that the users used to solve three workflow problems:

Problem #1: Off-Topic. The system occasionally failed to maintain a clear topic definition or emphasis over generations. To rectify this, users either reiterated the topic for relevance or added definitions and more details:

1a. Reiterate Topic (22 edits): When the system veered off the technical topic during a lengthy generation sequence, users reinforced the theme. They either reiterated their Tweakorial topic or inserted the reference ID, like [tweet_topic], into the prompt to keep generations relevant.

1a. Clarify Definition (5 edits): Occasionally, users perceived the system as lacking depth or contemporary knowledge about specific topics. To address this, they either embedded a direct definition or sourced definitions from online platforms like Wikipedia.

Problem #2: Undesired Narrative and Style. The system’s outputs sometimes didn’t align with users’ desired narrative or stylistic preferences. To personalize and guide the generations, users either add preferred examples as training to guide generations or provide longer instructions:

2a. Guide Narrative (25 edits): The system sometimes missed users’ preferred narratives or specific nuances they wanted to highlight. Users then reinforced their desired narrative direction by adding specifics or personal experiences.

2b. Guide Style (18 edits): To better mirror their unique writing styles, users made adjustments. Whether they wanted casual language, content tailored for a teenage audience, or the inclusion of emojis, users fine-tuned the prompt to achieve the desired output.

Problem #3: Inefficient Workflow. The system sometimes failed to generate users’ anticipated outcomes efficiently after familiarization. To utilize preferred, shorter version of the workflow or explore more options at once, they either sidestepped some middle steps or requested the system to diversify generations:

3a. Increase Variations (14 edits): Users observed that the original system prompt often generated a restricted number of examples. To expand their options, they modified the prompt to produce between 10 and 20 example outputs for broader brainstorming.

3b. Skip Steps (5 edits): When users identified valuable ideas early on, they sometimes wished to advance without adhering to every step. To streamline their process, they altered the system’s reference ID to reflect the results of the previous step.

Table 3: Overview of the 6 prompt editing strategies and their examples.

Problems	Editing Strategies	Example (Note: Highlights indicate user additions or deletions (striketrough) to the original prompt)
Problem 1: Off-Topic	Reiterate Topic (22 edits)	Please list 5 user scenarios someone encounters with [common_experience] and [tweet_topic] and augmented reality . Each scenario should be under 15 words, from a first-person perspective, related to a general audience, and not require technical knowledge. (P12, Step 3: Generate User Scenarios)
	Clarify Definition (5 edits)	Please list 5 everyday examples of [tweet_topic]. [tweet topic] is defined as using computational methods, often natural language processing, to analyze social science such as literature . Each example should be under 10 words, related to a general audience, and not require technical knowledge. (P4, Step 1: Generate Everyday Examples)
Problem 2: Undesired Narrative and Style	Guide Narrative (25 edits)	Here is a personal anecdote about the topic, [tweet_topic], for the experience of [common_experience]: [specific_anecdote]. Based on this anecdote, please write 3 versions of the opening tweet to introduce the concept of [tweet_topic]. The tweet should have no jargon. It should contain a specific and relatable example. It should have a driving question to spark curiosity and make readers want to read more. It should emphasize giving everyone an equal number of candies in broad categories, like chocolates, caramels, and that some candies can be in multiple categories . (P2, Step 6: Generate Example Hooks)
	Guide Style (18 edits)	Please list 5 common experiences with using [everyday_example]. Each experience should be under 10 words, related to a teenage audience , and not require technical knowledge. (P4, Step 2: Generate Common Experiences)
Problem 3: Inefficient Workflow	Increase Variation (14 edits)	Please list 5 20 user scenarios someone encounters with [common_experience]. Each scenario should be under 15 words, from a first-person perspective, related to a general audience, and not require technical knowledge. (P9, Step 3: Generate User Scenarios)
	Skip Step (5 edits)	Please write 3 short personal anecdotes about the topic, [tweet_topic], for the user scenario of [user_scenario] . Each anecdote should be under 30 words, from a first-person perspective, related to a general audience, and not require technical knowledge. (*Note: The user skipped three steps from Step #1 to #4, so they only used [tweet_topic] for this step, without chaining the [user_scenario].) (P7, Step 4: Generate Personal Anecdotes)

Augmenting Prompts with Users' Hooks

Starting from the 8th session (day one of week three), we introduced a new feature: the Learn-My-Style option inside 'Step 6: Generate Example Hooks'. This step allows the user to see two versions of the prompt: one is the original generate-example-hook prompt, and the other combines the original prompt with their own hook examples taken from their previous seven sessions. Users have the option to select one of the prompts to proceed, and they can also edit their own hook examples within the prompt if they wish (See the updated interface in Appendix B).

In theory, this Learn-My-Style prompt could enhance the system's generations by better adhering to users' writing styles. It also has the potential to provide more context about the topics or expertise areas users usually select from.

However, the session surveys revealed no statistically significant difference in users' perceived system usefulness and performance before and after we implemented the new feature. Specifically, between week two (sessions 5, 6, 7) and week three (sessions 8, 9, 10), there was no statistically significant difference in perceived usefulness for the overall system (6.03/7 - 6.06/7, p -value = 0.87) and for Step 6 (5.83/7 - 6.08/7, p -value = 0.47). Additionally, there was no statistically significant difference in performance (5.91/7 - 6.03/7, p -value = 0.40).

During the final interview, nine users expressed positive opinions on this feature, while three others showed either indifference or did not like the addition. All nine users who favored the new prompt compared the generated hooks with and without adding their own hook training data to understand how it works. Among them, four noted that their unique style was reflected in the generated hooks. For instance, P4 stated their preference for humor and darker themes, commenting that the hooks generated by the new prompt resonated well with their taste. P7 found the new prompt successfully captured their habit of ending tweets with a question and avoiding direct mention of topic names. P10 mentioned they like using personal pronouns, like "you," "she," "my," etc., to make hooks more relatable, and they saw this trait also from using the Learn-My-Style prompt.

Among the users less enthusiastic about this feature, P6 expressed hesitancy, feeling they needed more time for familiarization. P13 preferred using the original prompt, finding its simplicity and current results satisfactory. P9 offered a unique perspective: they saw the system as a tool to explore styles beyond their own. They believed the original system could offer a broader creative spectrum than their own style while incorporating their previous generations inside the prompt might constrain creativity and limit variety:

"I [already] had some biases in selecting hooks... [When I used] the original hook system, it might have just like more creative [and] expressive results. But then, I also tried to use the Learn-My-Style thing, and I was like, yeah, the creativity does get hindered."
— P9, Week 3 Interview

In conclusion, prompt editing becomes more useful over time. After users learn how and what they can add to prompts to enhance them, they report finding it significantly more valuable. Users were

divided on whether adding their hooks into the prompt as an example was beneficial. Even for those who found it useful, it addressed surface-level aspects like style rather than deeper structural issues like constraining the generations or skipping steps. Thus, AI generation systems would likely benefit users in the long term by allowing end users to edit prompts, enabling them to make structural changes to the system.

RQ 3: MENTAL MODELS AND OWNERSHIP

How have users' mental models, ownership, and involvement changed over time?

Summary: Users' mental models remain fairly consistent over time. However, users exhibit different levels of involvement and ownership throughout longitudinal interactions. Eight of 12 users thought that more involvement with the system led to increased feelings of ownership.

Mental Models

We analyzed the survey question on how users would describe the system throughout the 120 sessions. They primarily viewed the system as a tool (90.7% of the sessions, 107 times) and content generator (84.7% of the sessions, 100 times), with less frequent perceptions like co-pilot (44.1% of the sessions, 52 times), reference guide (38.1% of the sessions, 45 times), and ghostwriter (28.8% of the sessions, 34 times).

Before and after familiarization, there is no significant change in the user's mental model of the system (See Table 4). After familiarization, there was a small increase in references to the system as a "content generator" and "collaborator", with increases of 11.70% and 11.45% respectively. Descriptions of the system as a "tool" decreased by 7.27%. However, these shifts were subtle and not statistically significant. This indicates that users maintain a consistent perception and mental model of the system before and after becoming familiar with it. This suggests that when the novelty wears off, users don't usually process new mental models for such a customizable system.

Table 4: Users' mental model of the system before and after familiarization (percentage of the session where users described the system by its role, role names adapted from [80])

Description of the system	Before fam.	After fam.	Change
Tool	92.98%	85.71%	-7.27%
Content Generator	77.19%	88.89%	11.70%
Co-pilot	40.35%	46.03%	5.68%
Reference Guide	38.60%	36.51%	-2.09%
Collaborator	29.82%	41.27%	11.45%

Ownership and Involvement

After every interaction with the system, we asked users their degree of ownership over the hook³. On average, users perceive a 13.3%

³Survey question: Please rate the following statement from 1 - 7 (Strongly Disagree - Strongly Agree): "In this round, I feel a strong sense of ownership over the final hook I wrote, reflecting my personal style, effort, and creativity."

increase in ownership over hooks *after* familiarization (p-value < 0.005). Before familiarization, users had an average ownership perception of 4.28/7, showing their attitude towards ownership was neutral. After familiarization, their ownership perception increased to 5.21/7, showing they were between somewhat agree and agree that they felt strong ownership over their hooks. One reason the user P11 gave was that, over time, they put more effort into using the system and that their involvement contributed to increased feelings of ownership — from simple efforts like more carefully reading and selecting LLM outputs to making more edits to both the text and the prompts.

"Now thinking back, I would say I feel my ownership is increased... One [reason] is that I actually read through every sentence and every word to make sure I feel good about it... So just by reading through all of these [system generations and...] choosing the one that I agree with make me feel like I own it... Over the 3 weeks, I started to rewrite more of the hooks, and that actually makes me feel even stronger with my ownership."

— P11, Week 3 Interview

To investigate the relationship between involvement and ownership, we looked at our interview data. During our week three interviews, we asked the users, "How did your involvement within the sessions change after familiarization?" and "Do you find your level of involvement influences your sense of ownership?" From the interviews, there were four types of answers stemming from different experiences with the system:

1. More involved, more ownership: "My hard work pays off" (4 users). P6, P10, P11, and P12 shared that they became more involved in the process after familiarization. All of them mentioned that after gaining a better understanding of the system, they felt confident and comfortable enough to modify the prompt, thus dedicating more involvement and effort to the session. Specifically, P6, P10, and P12 iterated on their prompt edits multiple times to find a favorable output aligned with their interests. P11 and P12 focused on reading every result closely to select the best option. Additionally, P10, P11, and P12 frequently edited the outputs and conducted fact-checks through web searches to verify the information's accuracy. Thus, throughout this process, they all shared an increased sense of connection and ownership to the generated artifacts.

"After making more edits and being more engaged in the interaction, I feel a stronger connection to the tone and effort."

— P10, Week 3 Interview

2. Less involved, but maintain ownership. "I still lead the system" (2 users). P7 and P8 found themselves less involved over time but still maintained high ownership. They found the system reliable and capable enough to assist with the task after observing its commendable performance over a span of three weeks. Thus, they started to allow the system to take on more work and made fewer but more concise edits. They shared that they still *"led the*

system" (P8) and their guidance over the system still makes a huge impact on the final artifact.

"I feel like, even if I did no revision at all, I'd still feel a high ownership because I did all the selections [and] I did guide the system through... if I have specific revisions in the words or if I have a revision in the prompts, I may feel more ownership over it versus if I did nothing... but still it's overall pretty high ownership generally."

— P8, Week 3 Interview

3. Less involved, less ownership. "I am now just a reviewer" (4 users). P1, P3, P4, and P9 found themselves less involved with a drop in perceived ownership. Similar to P7 and P8, these four users also found the system reliable and capable. Moreover, P3 and P4 noted that they trusted the system more over time. Consequently, they started only providing high-level feedback to the system and reviewing its generation. For example, P3 expressed that, over usage, they increasingly allowed the system to handle more parts of the tasks. P3 mentioned they now acted more as reviewers:

"I know the system's capability better... I can trust its ability to generate some everyday examples and a final hook... I have become more reliant on the system... I become more like a human reviewer to help guide the GPT to generate a better hook... [so now] I think my ownership dropped a lot..."

— P3, Week 3 Interview

4. No changes in involvement/ownership over time. "Topic dictates effort" (2 users). P2 and P13 didn't report changes in their involvement or perception of ownership. They both found their experience was predominantly influenced by how well the LLM was able to support their chosen topic. P2 mentioned that the involvement increased only when they found the system failed to support their hook topics. P13 mentioned that they cannot fully trust the LLMs and have to intervene along the way. But they both agreed that the more edits they made, the more ownership they felt, but it wasn't a function of time. P13 compared their system experiences to mentoring undergraduate interns:

"It's like an undergrad intern... I had to teach [the system about my topic]. I couldn't fully trust GPT... [It's like] working with my undergrad, and I tell him what to code [...but not] how to code... Then he comes back to me, and we go over [the code to see...] if there's anything that needs more rigorous thinking... what's wrong... and how should we deal with it."

— P13, Week 3 Interview

In conclusion, although average feelings of ownership increased, there were nuanced ways that users experienced ownership after multiple experiences with the system. There does seem to be a relationship between effort and ownership. Based on the current state of LLM technology, many users find it necessary to edit text and prompts at every stage, but it is also a function of the topic - some topics are easier for LLMs to get right.

7 DISCUSSION

7.1 Why does the novelty not wear off?

This paper presents a three-week study of an LLM-based workflow to support science writing. Though generative AI has impressive abilities at first glance, it remains unclear whether those abilities stem from the initial impression of a new technology or from the genuine utility of the system's workflow. This study found that there was a familiarization phase (lasting 4.3 sessions on average) wherein users were exploring the system and its capabilities. After familiarization, they no longer explored its novel capabilities but focused solely on performing the task. They understood the workflow better, had a better ability to anticipate system generation, and performed better in writing hooks. Surprisingly, we found the average system utility *increased* by 12.1% after familiarization (5.47/7 to 6.32/7, p -value < 0.005), indicating that not only that the novelty does not wear off, but the workflow becomes more useful over time. We discuss potential reasons that the novelty did not wear off.

Reason #1: The workflow is supporting a task with cognitively challenging steps. One possible reason the workflow maintains its usefulness over time is that the underlying task targeted by the AI workflow has steps with a high cognitive load. A formative study with scientists writing Tweaktorials confirmed it was a cognitively challenging task [26]. When analyzing the cognitive load through the cognitive process theory of writing [23], they found there were two mentally demanding steps that writers wanted help with: *ideation* and *translation* (drafting complex thoughts into a linear narrative). This workflow helps with both these demanding steps, and they are the types of tasks that generally do not get easier over a short period of time.

The Tweaktorial Hooks workflow supports *ideation* in steps 1, 2, and 3. Together, these help users identify a specific and relatable use case that will be the vehicle for hooking non-technical audiences. Using the LLM to brainstorm multiple options helps users overcome fixation. Many workflows, especially generative AI-based ones, have shown that brainstorming support increases ideas' number, diversity, and quality [58, 59, 72, 91, 100, 112].

The Tweaktorial Hook workflow supports *translating* the use case into prose that is jargon-free, relatable, and sparks curiosity. Translation is difficult because it forces users to synthesize many pieces of information into one coherent, linear text. Often, the shorter the text, the harder it is to write, but generative AI can help produce examples of linear prose that satisfy some or all of these design goals. Thus, it can be an excellent starting point or motivation to overcome the mental activation energy needed to start drafting. Even when LLMs make mistakes, producing example text to work from helps users get over the cognitive barrier of translation. Many other writing support systems have shown that LLMs can assist with translation [52, 65, 74, 85, 87, 99, 109, 112].

In general, we expect that AI workflows that prove to be useful in a single session will likely maintain their usefulness over time if they assist with fundamentally challenging tasks. This includes cognitive challenges like *ideation*, sensemaking challenges like *clustering* and *synthesis*, socially challenging tasks like *providing multiple user perspectives*, and tedious or time-consuming tasks like *data formatting and processing*. Moreover, the

utility of AI workflows can increase over time as users become more efficient at using it, better at applying it to their domain, and can easily customize it to their needs.

Reason #2: Workflows improve human-AI collaboration and mitigate the effect of AI errors. Another possible reason that the workflow maintained its usefulness over time is that the task decomposition and planning provided by the workflow made it easier for users to solve the problem and to catch and correct errors made by the AI before they propagated. Often, decomposing a hard prompt into subtasks makes it easier to solve. This has been shown for an individual doing a task [7], for groups of people doing a task [6, 55, 78] and for LLMs doing tasks without humans [88, 101, 103]. Additionally, planning is also a mentally difficult process. Writers struggle to mentally shift between planning tasks and execution tasks, especially in writing [23]. Workflows that provide a plan and help people avoid task switching between high-level planning and low-level execution tasks are poised to bring lasting value.

Because workflows provide planning and task decomposition, it is fitting that workflows can enable better human-AI collaboration. In this AI workflow and many others [58, 72, 99, 112] AI is used to produce ideas, suggestions, drafts, and prototypes, but users control the outcome at every stage by judging the outputs and editing, regeneration or rewriting when the AI outputs are poor. Generative AI is known to make many mistakes. And there are tasks that generative AI does not know how to do. (For example, it is not particularly good at rating its own outputs.) Human judgment at every step ensures that AI errors do not propagate into bigger problems that are then harder to fix. The novelty of a workflow is unlikely to wear off if it genuinely supports the problem-solving process for the human. AI suggestions for each of the steps help the user. Since workflows limit the damage of AI errors and make them possible to correct before they propagate, users are likely to keep using the system rather than abandon it due to AI errors.

In this study, users showed a strong preference for using the workflow rather than trying to correct the hooks fully written by AI. In step 1 of the Tweaktorial Hook workflow, there is a "quick hook" where the LLM generates a hook for the topic in one prompt (without the workflow steps). Users were allowed to use the quick hook as their final hook and conclude the writing session if they wanted to. However, users rarely took it: out of the 120 sessions, users only used the quick hook 13 times. Most often, it was either factually incorrect or poorly constructed. Thus, many users expressed willingness to invest more effort in utilizing the entire workflow rather than try to fix a problematic quick hook. For example, one quick hook for the topic of "AJAX" based its hook on an incorrect usage of the technology. With an incorrect example at the core of the hook, it was impossible for the user to make simple edits to correct it; it is easier to start over and use a workflow, where users could select correct usages early and build toward a hook together with the AI. Additionally, by breaking the task into an AI workflow, there are more opportunities to customize the system at individual steps, for example, producing 20 user scenarios rather than 5, or specifying the common experiences should be relatable to a teenage audience.

In general, we expect that AI workflows will retain their value when the task decomposition offers problem-solving benefits to human users (even without AI). Moreover, the

workflow enables better human-AI collaboration by allowing users to check and edit AI outputs to control the quality of the output at every step.

7.2 Customization and the Potential of Appropriation in Generative AI Workflows

In this longitudinal study, we saw a significant increase in system utility due to a simple customization feature of exposing the prompt box so users could edit it. We believe this is a simple and powerful technology that all AI workflows should explore to help users mold the system to their needs. First, the prompt box helps users understand and explore the human-AI workflow rather than turning it into a “double black box” or hiding the affordances behind menus [62]. The prompt box facilitates transparency and offers a direct mental model of how the workflow operates in the background, enhancing user control and potential error recovery. Moreover, AI poses unique opportunities for end-user customization within a system. Traditional workflows are often brittle, constraining, and difficult for end users to edit [77]. However, with natural language editing, users can customize workflows in their own words.

Beyond customization, there is the potential for users to appropriate a workflow for tasks it was not originally designed for. In the seminal paper on Designing for Appropriation [16], Alan Dix summarizes 7 guidelines for building systems that “plan for the expected”. In the age of generative AI, many of these principles have become trivial to apply to systems. Systems can “provide visibility” by exposing prompts; they can “support not control” by allowing editing. The very nature of prompting with natural language “allows for interpretation (Do not make everything in the system or product have a fixed meaning, but include elements where users can add their own meanings).” If the system logs users’ prompt edits, it becomes possible to “learn from users’ appropriations” at an astonishing speed. At an extreme, it presents an opportunity for a continuous co-design process by sharing successful appropriations with the user community.

In our interviews, users saw the potential to appropriate the system. For example, P2 initially encountered difficulties when they found that their theoretical topics did not align well with the existing workflow, which primarily focused on everyday use cases. Thus, starting from session 4, they began to appropriate the system by using analogies rather than use cases to make theoretical topics relatable. After some prompt engineering, the metaphor strategy started to work. This could be a valuable approach for other theoretical topics that lack relatable use cases. Moreover, towards the end of the study, seven users shared expectations to appropriate the workflow to support other science communication tasks, like writing and revising their own research papers (P1, P8, P9, P10, P12, P13), drafting an elevator pitch about their research (P10), or making learning materials for themselves (P11). These are all interesting directions for this workflow to go in, and we are excited to see whether users can take it there on their own.

In general, future human-AI workflows should provide editable prompt boxes so users can customize workflows using natural language. This supports diverse user needs and has the potential to allow users to appropriate the system for tasks that it was not originally designed to support.

8 LIMITATIONS AND FUTURE WORK

Our study focused on a specific AI workflow utilizing a particular generative model. While we believe that this choice provides insights that can be generalized to other workflows, some findings are particular to the system studied here. The familiarization period in this study lasted 4.27 sessions on average. This is particular to the Tweakorial Hook system. A workflow that was more complex or a task that was more unfamiliar might require more time to explore. Particularly, workflows that are less linear and have more pathways will take more time for users to explore.

The Tweakorial Hook workflow uses GPT-4, but future LLMs might have better capabilities. If LLMs improve, will workflows still be needed? If LLMs become more factually correct, users may not need workflows to identify correct applications of a technology – they might be able to rely on LLMs for that. However, even if LLMs improve on factuality, they might not improve in their ability to know what examples will resonate with a particular audience, so human judgment in the workflow will likely be important even if LLM models improve.

Our study encompassed three weeks (ten usage sessions) and 11 of 12 users reported they already reached familiarization as they fully understood the technology. However, it is if users continued to use the workflow, their behavior would continue to change. First, it is possible that some aspects of the system would be less novel, if the user starts to notice more patterns in the LLM responses. However, even if the patterns become more clear, the system outputs could still be useful as a creative springboard. Second, we would hope to see more customization over time. As users repeat tasks, there are increased benefits to customization. Even a finalized workflow can continue to be iterated on for increased efficiency. We think there are many research questions that could be answered with more longitudinal studies of AI-based systems.

This study did not investigate the dangers of users potentially becoming dependent on an AI tool. Future studies could look at human performance on a task after the AI tools they used were taken away. This is important because some tools that students are allowed to use in academic settings might not be possible to use in corporate jobs due to data security concerns with AI systems. AI might be deskilling workers, but by providing scaffolding to complete a task, it also might help workers achieve skills. This deserves further study.

Our study relied largely on self-reporting to identify the familiarization period. The novelty effect is primarily a cognitive phenomenon associated with liking or pleasure, and thus, some amount of self-assessment is common to use [40]. However, more objective metrics might be applied in addition to self-assessment: measuring time spent on a task, measuring task performance (assessed through external evaluation), and measuring task competency through quizzes. Additionally, a study in the wild could use continued use or abandonment as a signal of a novelty effect.

Our participant pool was limited to Computer Science PhD students, a demographic proficient with technology. Even though they possess differing levels of experience with using generative AI – ranging from daily to never – their relatively high technology/AI literacy and openness to accepting new technology might differ from that of the general population. This could still introduce potential

bias. Additionally, users were not required to post their hooks publicly, which could have influenced their interaction with the system. If the stakes were higher, their critique of the system and session experiences might have been more stringent, or they might have invested more effort. Future studies might consider diversifying the user base or altering the stakes to observe any shifts in results.

9 BROADER IMPACT

Our study presents findings that could apply to both the general audience and academia: First, our study is among the earliest to examine the novelty effects and long-term utility of generative AI-based workflows. Thus, the study results shed light on insights for the future development of AI models and workflows based on them. Additionally, the study contributes a methodological foundation for future empirical studies, particularly long-term investigations, on generative AI and its workflows. Second, many of the insights align well with existing design theories, such as technology domestication theory, technology acceptance model, and human-AI interaction guidelines, thereby offering new insights for designing longitudinal interactions in the era of generative AI and AI-powered creativity support tools. Third, our paper highlights the potential of technology appropriation, which allows and encourages users to customize and adapt systems for new use cases. We initiate a discussion about using generative AI as design material to promote user-centric co-design for system development.

However, we also acknowledge that our study primarily reports positive user interaction with AI workflows over time, lacking emphasis on the potential risks and harms associated. Our user group might have already encountered these risks and harms in various contexts before the study. This familiarity can breed trust and contribute to higher tolerance levels, potentially leading to the oversight or disregard of associated harms observed in the study.

10 CONCLUSION

To understand the novelty and long-term utility of AI workflows, we conducted a three-week longitudinal study with 12 users to understand the customization of generative AI workflow for science communication. Our study revealed that the familiarization phase lasts for 4.27 sessions, which equates to 1.42 hours or one and a half weeks of usage. Before familiarization, users are mainly exploring the capabilities of the workflow and which aspects they find useful. After familiarization, the perceived utility of the system is rated 12.1% higher than before, indicating that the perceived utility of AI is not just a novelty effect. The increase in benefits mainly comes from end-users' ability to customize prompts, and thus appropriate the system to their own needs, ranging from addressing system problems to adapting for new use cases. Though the mental models remain over time, users begin to change their level of involvement while collaborating with the AI workflow, thus leading to a shift in their perception of ownership. This points to a future in which generative AI systems and workflows can facilitate user customization and appropriation, as well as a continuous co-design process between system developers and users.

11 ACKNOWLEDGEMENT

This work is supported by NSF IIS-2202578 and NSF IIS-2129020. We thank Dorothy Zhang, Karthik Sreedhar, and Weirui Peng for their help in setting up the study. We thank Sitong Wang, Vivian Liu, Soya Park, Chinmay Kulkarni, the Interactive Creativity Lab at Adobe Research, and the anonymous reviewers for their feedback on the draft.

REFERENCES

- [1] Ian Arawjo, Chelse Swoopes, Priyan Vaithilingam, Martin Wattenberg, and Elena Glassman. 2023. ChainForge: A Visual Toolkit for Prompt Engineering and LLM Hypothesis Testing. <https://doi.org/10.48550/arXiv.2309.09128>
- [2] Roy F Baumeister and Mark R Leary. 1995. The need to belong: desire for interpersonal attachments as a fundamental human motivation. *Psychological bulletin* 117, 3 (1995), 497.
- [3] Amaury Belin and Yannick Prié. 2012. DIAM: Towards a model for describing appropriation processes through the evolution of digital artifacts. In *Proceedings of the Designing Interactive Systems Conference* (Newcastle Upon Tyne, United Kingdom) (DIS '12). Association for Computing Machinery, New York, NY, USA, 645–654. <https://doi.org/10.1145/2317956.2318053>
- [4] Ivo Benke, Maren Schneider, Xuanhui Liu, and Alexander Maedche. 2022. Team-Spiritous - A Retrospective Emotional Competence Development System for Video-Meetings. *Proceedings of the ACM on Human-Computer Interaction* 6, CSCW2 (Nov 2022), 1–28. <https://doi.org/10.1145/3555117>
- [5] Thomas Berker, Maren Hartmann, and Yves Punie. 2005. *Domestication of media and technology*. McGraw-Hill Education (UK).
- [6] Michael S. Bernstein, Greg Little, Robert C. Miller, Björn Hartmann, Mark S. Ackerman, David R. Karger, David Crowell, and Katrina Panovich. 2010. Soylent: a word processor with a crowd inside. In *Proceedings of the 23rd Annual ACM Symposium on User Interface Software and Technology* (New York, New York, USA) (UIST '10). Association for Computing Machinery, New York, NY, USA, 313–322. <https://doi.org/10.1145/1866029.1866078>
- [7] Ananya Bhattacharjee, Yuchen Zeng, Sarah Yi Xu, Dana Kulzhabayeva, Minyi Ma, Rachel Kornfield, Syed Ishtiaque Ahmed, Alex Mariakakis, Mary P Czerwinski, Anastasia Kuzminykh, et al. 2023. Understanding the Role of Large Language Models in Personalizing and Scaffolding Strategies to Combat Academic Procrastination. *arXiv preprint arXiv:2312.13581* (2023).
- [8] Jessica Blaynee, Dr David Kreps, Dr Maria Kutar, and Dr Marie Griffiths. 2016. Collaborative HCI and UX: longitudinal diary studies as a means of uncovering barriers to digital adoption. In *Proceedings of the 30th International BCS Human Computer Interaction Conference: Fusion!* (Poole, United Kingdom) (HCI '16). BCS Learning & Development Ltd., Swindon, GBR, Article 53, 6 pages. <https://doi.org/10.14236/ewic/HCI2016.72>
- [9] Rishi Bommasani, Drew A. Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S. Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, Erik Brynjolfsson, Shyamal Buch, Dallas Card, Rodrigo Castellon, Niladri Chatterji, Annie Chen, Kathleen Creel, Jared Quincy Davis, Dora Demszky, Chris Donahue, Moussa Doumbouya, Esin Durmus, Stefano Ermon, John Etchemendy, Kavin Ethayarajh, Li Fei-Fei, Chelsea Finn, Trevor Gale, Lauren Gillespie, Karan Goel, Noah Goodman, Shelby Grossman, Neel Guha, Tatsunori Hashimoto, Peter Henderson, John Hewitt, Daniel E. Ho, Jenny Hong, Kyle Hsu, Jing Huang, Thomas Icard, Saahil Jain, Dan Jurafsky, Pratyusha Kalluri, Siddharth Karamcheti, Geoff Keeling, Fereshte Khani, Omar Khattab, Pang Wei Koh, Mark Krass, Ranjay Krishna, Rohith Kudithipudi, Ananya Kumar, Faisal Ladhak, Mina Lee, Tony Lee, Jure Leskovec, Isabelle Levent, Xiang Lisa Li, Xuechen Li, Tengyu Ma, Ali Malik, Christopher D. Manning, Suvir Mirchandani, Eric Mitchell, Zanele Muniyikwa, Suraj Nair, Avani Narayan, Deepak Narayanan, Ben Newman, Allen Nie, Juan Carlos Nieves, Hamed Nilforoshan, Julian Nyarko, Giray Ogut, Laurel Orr, Isabel Papadimitriou, Joon Sung Park, Chris Piech, Eva Portelance, Christopher Potts, Aditi Raghunathan, Rob Reich, Hongyu Ren, Frieda Rong, Yusuf Roohani, Camilo Ruiz, Jack Ryan, Christopher Ré, Dorsa Sadigh, Shiori Sagawa, Keshav Santhanam, Andy Shih, Krishnan Srinivasan, Alex Tamkin, Rohan Taori, Armin W. Thomas, Florian Tramèr, Rose E. Wang, William Wang, Bohan Wu, Jiajun Wu, Yuhuai Wu, Sang Michael Xie, Michihiro Yasunaga, Jiaxuan You, Matei Zaharia, Michael Zhang, Tianyi Zhang, Xikun Zhang, Yuhui Zhang, Lucia Zheng, Kaitlyn Zhou, and Percy Liang. 2022. On the Opportunities and Risks of Foundation Models. *arXiv:2108.07258* [cs] <http://arxiv.org/abs/2108.07258>
- [10] Massimiano Bucchi and Brian Trench. 2021. *Routledge Handbook of Public Communication of Science and Technology*. Routledge. Google-Books-ID: Em4ZEAAAQBAJ.
- [11] Victor Bursztyn, David Demeter, Doug Downey, and Larry Birnbaum. 2022. Learning to Perform Complex Tasks through Compositional Fine-Tuning of Language Models. In *Findings of the Association for Computational Linguistics*:

- EMNLP 2022 (Abu Dhabi, United Arab Emirates, 2022-12), Yoav Goldberg, Zornitsa Kozareva, and Yue Zhang (Eds.). Association for Computational Linguistics, 1676–1686. <https://doi.org/10.18653/v1/2022.findings-emnlp.121>
- [12] Jennie Carroll, Steve Howard, Frank Vetere, Jane Peck, and John Murphy. 2001. Identity, Power And Fragmentation in Cyberspace: Technology Appropriation by Young People. (01 2001).
- [13] DaEun Choi, Sumin Hong, Jeongeun Park, John Joon Young Chung, and Juho Kim. 2023. CreativeConnect: Supporting Reference Recombination for Graphic Design Ideation with Generative AI. *arXiv:2312.11949* [cs] <http://arxiv.org/abs/2312.11949>
- [14] Nicholas Davis, Holger Winnemöller, Mira Dontcheva, and Ellen Yi-Luen Do. 2013. Toward a cognitive theory of creativity support. In *Proceedings of the 9th ACM Conference on Creativity & Cognition* (Sydney, Australia) (C&C '13). Association for Computing Machinery, New York, NY, USA, 13–22. <https://doi.org/10.1145/2466627.2466655>
- [15] Edward L Deci and Richard M Ryan. 2000. The “what” and “why” of goal pursuits: Human needs and the self-determination of behavior. *Psychological inquiry* 11, 4 (2000), 227–268.
- [16] Alan Dix. 2007. Designing for appropriation. In *Proceedings of the 21st British HCI Group Annual Conference on People and Computers: HCI...but Not as We Know It - Volume 2* (University of Lancaster, United Kingdom) (BCS-HCI '07). BCS Learning & Development Ltd., Swindon, GBR, 27–30.
- [17] Paul Dourish. 2003. The Appropriation of Interactive Technologies: Some Lessons from Placeless Documents. *Computer Supported Cooperative Work (CSCW)* 12, 4 (2003), 465–490. <https://doi.org/10.1023/A:1026149119426>
- [18] Yilun Du, Shuang Li, Antonio Torralba, Joshua B. Tenenbaum, and Igor Moradach. 2023. Improving Factuality and Reasoning in Language Models through Multiagent Debate. *arXiv:2305.14325* [cs] <http://arxiv.org/abs/2305.14325>
- [19] Jacob Eisenstein, Daniel Andor, Bernd Bohnet, Michael Collins, and David Minno. 2022. Honest Students from Untrusted Teachers: Learning an Interpretable Question-Answering Pipeline from a Pretrained Language Model. *arXiv:2210.02498* [cs] <http://arxiv.org/abs/2210.02498>
- [20] Tyna Eloundou, Sam Manning, Pamela Mishkin, and Daniel Rock. 2023. GPTs are GPTs: An Early Look at the Labor Market Impact Potential of Large Language Models. *arXiv:2303.10130* (Aug 2023). <http://arxiv.org/abs/2303.10130> *arXiv:2303.10130* [cs, econ, q-fin].
- [21] K. J. Kevin Feng, Q. Vera Liao, Ziang Xiao, Jennifer Wortman Vaughan, Amy X. Zhang, and David W. McDonald. 2024. Canvil: Designerly Adaptation for LLM-Powered User Experiences. *arXiv:2401.09051* [cs] <http://arxiv.org/abs/2401.09051>
- [22] Jennifer Fereday and Eimear Muir-Cochrane. 2006. Demonstrating rigor using thematic analysis: A hybrid approach of inductive and deductive coding and theme development. *International journal of qualitative methods* 5, 1 (2006), 80–92.
- [23] Linda Flower and John R. Hayes. 1981. A Cognitive Process Theory of Writing. *College Composition and Communication* 32, 4 (1981), 365–387. <https://eric.ed.gov/?id=EJ256235>
- [24] Jie Gao, Yuchen Guo, Gionnieve Lim, Tianqin Zhang, Zheng Zhang, Toby Jia-Jun Li, and Simon Tangi Perrault. 2024. CollabCoder: A Lower-barrier, Rigorous Workflow for Inductive Collaborative Qualitative Analysis with Large Language Models. *arXiv:2304.07366* [cs] <http://arxiv.org/abs/2304.07366>
- [25] Katy Ilonka Gero, Vivian Liu, and Lydia Chilton. 2022. Sparks: Inspiration for Science Writing Using Language Models. In *Designing Interactive Systems Conference* (Virtual Event, Australia) (DIS '22). Association for Computing Machinery, New York, NY, USA, 1002–1019. <https://doi.org/10.1145/3532106.3533533>
- [26] Katy Ilonka Gero, Vivian Liu, Sarah Huang, Jennifer Lee, and Lydia B. Chilton. 2021. What Makes Tweetorials Tick: How Experts Communicate Complex Topics on Twitter. *Proc. ACM Hum.-Comput. Interact.* 5, CSCW2, Article 422 (oct 2021), 26 pages. <https://doi.org/10.1145/3479566>
- [27] Katy Ilonka Gero, Tao Long, and Lydia B Chilton. 2023. Social Dynamics of AI Support in Creative Writing. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (Hamburg, Germany) (CHI '23). Association for Computing Machinery, New York, NY, USA, Article 245, 15 pages. <https://doi.org/10.1145/3544548.3580782>
- [28] Saul Greenberg and Bill Buxton. 2008. Usability evaluation considered harmful (some of the time). In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Florence, Italy) (CHI '08). Association for Computing Machinery, New York, NY, USA, 111–120. <https://doi.org/10.1145/1357054.1357074>
- [29] Madeleine Grunde-McLaughlin, Michelle S. Lam, Ranjay Krishna, Daniel S. Weld, and Jeffrey Heer. 2023. Designing LLM Chains by Adapting Techniques from Crowdsourcing Workflows. *arXiv:2312.11681* [cs] <http://arxiv.org/abs/2312.11681>
- [30] Guhyun Han, Jaehun Jung, Young-Ho Kim, and Jinwook Seo. 2023. Data-Halo: A Customizable Notification Visualization System for Personalized and Longitudinal Interactions. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. ACM, Hamburg Germany, 1–21. <https://doi.org/10.1145/3544548.3580828>
- [31] Eszter Hargittai, Tobias Fuchsli, and Mike S. Schäfer. 2018. How Do Young Adults Engage With Science and Research on Social Media? Some Preliminary Findings and an Agenda for Future Research. *Social Media & Society* 4, 3 (July 2018), 2056305118797720. <https://doi.org/10.1177/2056305118797720> Publisher: SAGE Publications Ltd.
- [32] Sandra G Hart and Lowell E Staveland. 1988. Development of NASA-TLX (Task Load Index): Results of empirical and theoretical research. *Human mental workload* 1, 3 (1988), 139–183.
- [33] Md Tanim Hasan, Shaila Zaman, Amanveer Wesley, Panagiotis Tsiamyrtzis, and Ioannis Pavlidis. 2023. Sympathetic Activation in Deadlines of Deskbound Research - A Study in the Wild. In *Extended Abstracts of the 2023 CHI Conference on Human Factors in Computing Systems*. ACM, Hamburg Germany, 1–8. <https://doi.org/10.1145/3544549.3585585>
- [34] Paul Hekkert, Dirk Snelders, and Piet C. W. Van Wieringen. 2003. ‘Most advanced, yet acceptable’: Typicality and novelty as joint predictors of aesthetic preference in industrial design. *British Journal of Psychology* 94, 1 (2003), 111–124. <https://doi.org/10.1348/000712603762842147> _eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1348/000712603762842147>
- [35] Ari Holtzman, Jan Buys, Li Du, Maxwell Forbes, and Yeyin Choi. 2020. The Curious Case of Neural Text Degeneration. *arXiv:1904.09751* [cs] (Feb. 2020). <http://arxiv.org/abs/1904.09751>
- [36] Qirui Huang, Min Lu, Joel Lanir, Dani Lischinski, Daniel Cohen-Or, and Hui Huang. 2024. GraphiMind: LLM-centric Interface for Information Graphics Design. *arXiv:2401.13245* [cs] <http://arxiv.org/abs/2401.13245>
- [37] Hilary Hutchinson, Wendy Mackay, Bo Westerlund, Benjamin B. Bederson, Allison Druin, Catherine Plaisant, Michel Beaudouin-Lafon, Stéphane Conversy, Helen Evans, Heiko Hansen, Nicolas Roussel, and Björn Eiderbäck. 2003. Technology probes: inspiring design for and with families. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Ft. Lauderdale, Florida, USA) (CHI '03). Association for Computing Machinery, New York, NY, USA, 17–24. <https://doi.org/10.1145/642611.642616>
- [38] Daphne Ippolito, Reno Kriz, Maria Kustikova, João Sedoc, and Chris Callison-Burch. 2019. Comparison of Diverse Decoding Methods from Conditional Language Models. *arXiv:1906.06362* [cs] (June 2019). <http://arxiv.org/abs/1906.06362> *arXiv:1906.06362*
- [39] Mohit Jain and Ravin Balakrishnan. 2012. User learning and performance with bezel menus. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. 2221–2230.
- [40] Lucas M. Jenö, Vigdis Vandvik, Sigrunn Eliassen, and John-Arvid Grytnes. 2019. Testing the novelty effect of an m-learning tool on internalization and achievement: A Self-Determination Theory approach. *Computers & Education* 128 (2019), 398–413. <https://doi.org/10.1016/j.compedu.2018.10.008>
- [41] Peiling Jiang, Jude Rayan, Steven P. Dow, and Haijun Xia. 2023. Graphologue: Exploring Large Language Model Responses with Interactive Diagrams. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology* (San Francisco, CA, USA) (UIST '23). Association for Computing Machinery, New York, NY, USA, Article 3, 20 pages. <https://doi.org/10.1145/3586183.3606737>
- [42] Majeed Kazemitabaar, Runlong Ye, Xiaoning Wang, Austin Z. Henley, Paul Denny, Michelle Craig, and Tovi Grossman. 2024. CodeAid: Evaluating a Classroom Deployment of an LLM-based Programming Assistant that Balances Student and Educator Needs. *arXiv:2401.11314* [cs] <http://arxiv.org/abs/2401.11314>
- [43] Jeongyeon Kim, Sangho Suh, Lydia B Chilton, and Haijun Xia. 2023. Metaphorian: Leveraging Large Language Models to Support Extended Metaphor Creation for Science Writing. In *Proceedings of the 2023 ACM Designing Interactive Systems Conference* (Pittsburgh PA USA, 2023-07-10). ACM, 115–135. <https://doi.org/10.1145/3563657.3595996>
- [44] Tae Soo Kim, DaEun Choi, Yoonseo Choi, and Juho Kim. 2022. Stylette: Styling the Web with Natural Language. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (New Orleans, LA, USA) (CHI '22). Association for Computing Machinery, New York, NY, USA, Article 5, 17 pages. <https://doi.org/10.1145/3491102.3501931>
- [45] Tae Soo Kim, Yoonjoo Lee, Minsuk Chang, and Juho Kim. 2023. Cells, Generators, and Lenses: Design Framework for Object-Oriented Interaction with Large Language Models. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology* (San Francisco, CA, USA) (UIST '23). Association for Computing Machinery, New York, NY, USA, Article 4, 18 pages. <https://doi.org/10.1145/3586183.3606833>
- [46] Maria Kjærup, Mikael B. Skov, Peter Axel Nielsen, Jesper Kjeldskov, Jens Gerken, and Harald Reiterer. 2021. *Longitudinal Studies in HCI Research: A Review of CHI Publications From 1982–2019*. Springer International Publishing, Cham, 11–39. https://doi.org/10.1007/978-3-030-67322-2_2
- [47] Michael Koch, Kai von Luck, Jan Schwarzer, and Susanne Draheim. 2018. The novelty effect in large display deployments—Experiences and lessons-learned for evaluating prototypes. In *Proceedings of 16th European conference on computer-supported cooperative work-exploratory papers*. European Society for Socially Embedded Technologies (EUSSET).

- [48] Thomas Kosch, Robin Welsch, Lewis Chuang, and Albrecht Schmidt. 2023. The Placebo Effect of Artificial Intelligence in Human–Computer Interaction. *ACM Trans. Comput.-Hum. Interact.* 29, 6, Article 56 (jan 2023), 32 pages. <https://doi.org/10.1145/3529225>
- [49] Alina Kriskchowsky, Manfred Tscheligi, Katja Neureiter, Michael Muller, Anna Maria Polli, and Nervo Verdezoto. 2015. Experiences of Technology Appropriation: Unanticipated Users, Usage, Circumstances, and Design. (2015).
- [50] Poey Chin Lai. 2017. The literature review of technology adoption models and theories for the novelty technology. *JISTEM-Journal of Information Systems and Technology Management* 14 (2017), 21–38.
- [51] Ana Levordashka, Jamie Eastman, Eleni Anna Skoularik, Anca Salagean, Darren Cosker, and Danaë Stanton Fraser. 2023. An Exploration of Theatre Rehearsals in Social Virtual Reality. In *Extended Abstracts of the 2023 CHI Conference on Human Factors in Computing Systems*. ACM, Hamburg Germany, 1–7. <https://doi.org/10.1145/3544549.3585685>
- [52] Cheng Li, Mingyang Zhang, Qiaozhu Mei, Yaqing Wang, Spurthi Amba Hombaiah, Yi Liang, and Michael Bendersky. 2023. Teach LLMs to Personalize – An Approach inspired by Writing Education. arXiv:2308.07968 [cs] <http://arxiv.org/abs/2308.07968>
- [53] Grace Li, Tao Long, and Lydia B Chilton. 2023. Eliciting Topic Hierarchies from Large Language Models. arXiv preprint arXiv:2310.19275 arXiv:2310.19275 (2023). <http://arxiv.org/abs/2310.19275>
- [54] Joseph Lindley, Paul Coulton, and Miriam Sturdee. 2017. Implications for Adoption. In *Proceedings of the 2017 CHI Conference on Human Factors in Computing Systems (CHI '17)*. Association for Computing Machinery, New York, NY, USA, 265–277. <https://doi.org/10.1145/3025453.3025742>
- [55] Greg Little, Lydia B. Chilton, Max Goldman, and Robert C. Miller. 2010. TurkKit: human computation algorithms on mechanical turk. In *Proceedings of the 23rd Annual ACM Symposium on User Interface Software and Technology* (New York, New York, USA) (UIST '10). Association for Computing Machinery, New York, NY, USA, 57–66. <https://doi.org/10.1145/1866029.1866040>
- [56] Vivian Liu and Lydia B Chilton. 2022. Design Guidelines for Prompt Engineering Text-to-Image Generative Models. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (New Orleans, LA, USA) (CHI '22). Association for Computing Machinery, New York, NY, USA, Article 384, 23 pages. <https://doi.org/10.1145/3491102.3501825>
- [57] Vivian Liu, Tao Long, Nathan Raw, and Lydia Chilton. 2023. Generative Disco: Text-to-Video Generation for Music Visualization. arXiv:2304.08551 (Apr 2023). <https://doi.org/10.48550/arXiv.2304.08551> arXiv:2304.08551 [cs].
- [58] Vivian Liu, Han Qiao, and Lydia Chilton. 2022. Opal: Multimodal Image Generation for News Illustration. In *Proceedings of the 35th Annual ACM Symposium on User Interface Software and Technology* (Bend, OR, USA) (UIST '22). Association for Computing Machinery, New York, NY, USA, Article 73, 17 pages. <https://doi.org/10.1145/3526113.3545621>
- [59] Vivian Liu, Jo Vermeulen, George Fitzmaurice, and Justin Matejka. 2023. 3DALL-E: Integrating Text-to-Image AI in 3D Design Workflows. In *Proceedings of the 2023 ACM Designing Interactive Systems Conference* (<conf-loc>, <city>Pittsburgh</city>, <state>PA</state>, <country>USA</country>, </conf-loc>) (DIS '23). Association for Computing Machinery, New York, NY, USA, 1955–1977. <https://doi.org/10.1145/3563657.3596098>
- [60] Tao Long and Lydia B. Chilton. 2023. Challenges and Opportunities for the Design of Smart Speakers. <https://doi.org/10.48550/arXiv.2306.05741> arXiv:2306.05741 [cs]
- [61] Tao Long, Dorothy Zhang, Grace Li, Batool Taraf, Samia Menon, Kynneddy Simone Smith, Sitong Wang, Katy Ilonka Gero, and Lydia B. Chilton. 2023. Tweetorial Hooks: Generative AI Tools to Motivate Science on Social Media. In *Proceedings of the 14th Conference on Computational Creativity (ICCC '23)*. Association for Computational Creativity, Ontario, Canada, 44–53.
- [62] Wendy E. Mackay. 1991. Triggers and barriers to customizing software. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (New Orleans, Louisiana, USA) (CHI '91). Association for Computing Machinery, New York, NY, USA, 153–160. <https://doi.org/10.1145/108844.108867>
- [63] Allan MacLean, Kathleen Carter, Lennart Löfstrand, and Thomas Moran. 1990. User-tailorable systems: pressing the issues with buttons. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Seattle, Washington, USA) (CHI '90). Association for Computing Machinery, New York, NY, USA, 175–182. <https://doi.org/10.1145/97243.97271>
- [64] Sampada Marathe and S. Shyam Sundar. 2011. What drives customization? control or identity?. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (Vancouver, BC, Canada) (CHI '11). Association for Computing Machinery, New York, NY, USA, 781–790. <https://doi.org/10.1145/1978942.1979056>
- [65] Piotr Mirowski, Kory W. Mathewson, Jaylen Pittman, and Richard Evans. 2023. Co-Writing Screenplays and Theatre Scripts with Language Models: Evaluation by Industry Professionals. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (Hamburg, Germany) (CHI '23). Association for Computing Machinery, New York, NY, USA, Article 355, 34 pages. <https://doi.org/10.1145/3544548.3581225>
- [66] Thomas P Moran. 2002. Everyday adaptive design. In *Proceedings of the 4th conference on Designing interactive systems: processes, practices, methods, and techniques* (London England, 2002-06-25). ACM, 13–14. <https://doi.org/10.1145/778712.778715>
- [67] Anirban Mukherjee and Hannah Chang. 2023. The Creative Frontier of Generative AI: Managing the Novelty-Usefulness Tradeoff. arXiv:2306.03601 [cs] <http://arxiv.org/abs/2306.03601>
- [68] Michael Muller, Katja Neureiter, Nervo Verdezoto, Alina Kriskchowsky, Anna Maria Al Zubaidi-Polli, and Manfred Tscheligi. 2016. Collaborative Appropriation: How Couples, Teams, Groups and Communities Adapt and Adopt Technologies. In *Proceedings of the 19th ACM Conference on Computer Supported Cooperative Work and Social Computing Companion* (San Francisco, California, USA) (CSCW '16 Companion). Association for Computing Machinery, New York, NY, USA, 473–480. <https://doi.org/10.1145/2818052.2855508>
- [69] Peter Axel Nielsen. 2021. Longitudinal Studies in Information Systems. In *Advances in Longitudinal HCI Research*, Evangelos Karapanos, Jens Gerken, Jesper Kjeldskov, and Mikael B. Skov (Eds.). Springer International Publishing, 41–56. https://doi.org/10.1007/978-3-030-67322-2_3
- [70] OpenAI. 2024. GPT-4 Technical Report. <https://doi.org/10.48550/arXiv.2303.08774> arXiv:2303.08774 [cs]
- [71] Joon Sung Park, Joseph O'Brien, Carrie Jun Cai, Meredith Ringel Morris, Percy Liang, and Michael S. Bernstein. 2023. Generative Agents: Interactive Simulacra of Human Behavior. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology* (San Francisco, CA, USA) (UIST '23). Association for Computing Machinery, New York, NY, USA, Article 2, 22 pages. <https://doi.org/10.1145/3586183.3606763>
- [72] Savvas Petridis, Nicholas Diakopoulos, Kevin Crowston, Mark Hansen, Keren Henderson, Stan Jastrzebski, Jeffrey V Nickerson, and Lydia B Chilton. 2023. AngleKindling: Supporting Journalistic Angle Ideation with Large Language Models. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (<conf-loc>, <city>Hamburg</city>, <country>Germany</country>, </conf-loc>) (CHI '23). Association for Computing Machinery, New York, NY, USA, Article 225, 16 pages. <https://doi.org/10.1145/3544548.3580907>
- [73] Savvas Petridis, Michael Terry, and Carrie Jun Cai. 2023. PromptInfuser: Bringing User Interface Mock-Ups to Life with Large Language Models. In *Extended Abstracts of the 2023 CHI Conference on Human Factors in Computing Systems* (Hamburg, Germany) (CHI EA '23). Association for Computing Machinery, New York, NY, USA, Article 237, 6 pages. <https://doi.org/10.1145/3544549.3585628>
- [74] Savvas Petridis, Ben Wedin, James Wexler, Aaron Donsbach, Mahima Pushkarna, Nitesh Goyal, Carrie J. Cai, and Michael Terry. 2023. ConstitutionMaker: Interactively Critiquing Large Language Models by Converting Feedback into Principles. arXiv:2310.15428 [cs] <http://arxiv.org/abs/2310.15428>
- [75] Mailson de Queiroz Proença, Vivian Genaro Motti, Kamila Rios da Hora Rodrigues, and Vânia Paula de Almeida Neris. 2021. Coping with Diversity - A System for End-users to Customize Web User Interfaces. *Proc. ACM Hum.-Comput. Interact.* 5, EICS, Article 201 (may 2021), 27 pages. <https://doi.org/10.1145/3457151>
- [76] Huilian Sophie Qiu, Anna Lieb, Jennifer Chou, Megan Carneal, Jasmine Mok, Emily Amspoker, Bogdan Vasilescu, and Laura Dabbish. 2023. Climate Coach: A Dashboard for Open-Source Maintainers to Overview Community Dynamics. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. ACM, Hamburg Germany, 1–18. <https://doi.org/10.1145/3544548.3581317>
- [77] Daniela Retelny, Michael S. Bernstein, and Melissa A. Valentine. 2017. No Workflow Can Ever Be Enough: How Crowdsourcing Workflows Constrain Complex Work. *Proc. ACM Hum.-Comput. Interact.* 1, CSCW, Article 89 (dec 2017), 23 pages. <https://doi.org/10.1145/3134724>
- [78] Daniela Retelny, Sebastian Robaszkiewicz, Alexandra To, Walter S. Lasecki, Jay Patel, Negar Rahmati, Tulsee Doshi, Melissa Valentine, and Michael S. Bernstein. 2014. Expert crowdsourcing with flash teams. In *Proceedings of the 27th annual ACM symposium on User interface software and technology - UIST '14*. ACM, New York, NY, USA, 75–85. <https://doi.org/10.1145/2642918.2647409>
- [79] Luiz Rodrigues, Filipe D. Pereira, Armando M. Toda, Paula T. Palomino, Marcela Pessoa, Leandro Silva Galvão Carvalho, David Fernandes, Elaine H. T. Oliveira, Alexandra I. Cristea, and Seiji Isotani. 2022. Gamification suffers from the novelty effect but benefits from the familiarization effect: Findings from a longitudinal study. *International Journal of Educational Technology in Higher Education* 19, 1 (2022). <https://doi.org/10.1186/s41239-021-00314-6> Publisher: SpringerOpen.
- [80] Steven I. Ross, Fernando Martinez, Stephanie Houde, Michael Muller, and Justin D. Weisz. 2023. The Programmer's Assistant: Conversational Interaction with a Large Language Model for Software Development. In *Proceedings of the 28th International Conference on Intelligent User Interfaces* (Sydney, NSW, Australia) (IUI '23). Association for Computing Machinery, New York, NY, USA, 491–514. <https://doi.org/10.1145/3581641.3584037>
- [81] Isa Rutten, Lawrence Van den Bogaert, and David Geerts. 2021. From Initial Encounter With Mid-Air Haptic Feedback to Repeated Use: The Role of the Novelty Effect in User Experience. *IEEE Transactions on Haptics* 14, 3 (2021), 591–602. <https://doi.org/10.1109/TOH.2020.3043658>

- [82] Isa Rutten and David Geerts. 2020. Better Because It's New: The Impact of Perceived Novelty on the Added Value of Mid-Air Haptic Feedback. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (CHI '20). Association for Computing Machinery, New York, NY, USA, 1–13. <https://doi.org/10.1145/3313831.3376668>
- [83] Richard M Ryan and Edward L Deci. 2017. *Self-determination theory: Basic psychological needs in motivation, development, and wellness*. Guilford publications.
- [84] Soheil Sadeghi, Somit Gupta, Stefan Gramatovici, Jiannan Lu, Hao Ai, and Ruhan Zhang. 2022. Novelty and Primacy: A Long-Term Estimator for Online Experiments. *Technometrics* 64, 4 (2022), 524–534. <https://doi.org/10.1080/00401706.2022.2124309> arXiv:https://doi.org/10.1080/00401706.2022.2124309
- [85] Timo Schick, Jane Dwivedi-Yu, Zhengbao Jiang, Fabio Petroni, Patrick Lewis, Gautier Izacard, Qingfei You, Christoforos Nalmpantis, Edouard Grave, and Sebastian Riedel. 2022. PEER: A Collaborative Language Model. arXiv:2208.11663 [cs] <http://arxiv.org/abs/2208.11663>
- [86] Woosuk Seo, Chanmo Yang, and Young-Ho Kim. 2024. ChaCha: Leveraging Large Language Models to Prompt Children to Share Their Emotions about Personal Events. arXiv:2309.12244 [cs] <http://arxiv.org/abs/2309.12244>
- [87] Hua Shen, Chieh-Yang Huang, Tongshuang Wu, and Ting-Hao Kenneth Huang. 2023. ConvXAI: Delivering Heterogeneous AI Explanations via Conversations to Support Human-AI Scientific Writing. In *Companion Publication of the 2023 Conference on Computer Supported Cooperative Work and Social Computing* (Minneapolis, MN, USA) (CSCW '23 Companion). Association for Computing Machinery, New York, NY, USA, 384–387. <https://doi.org/10.1145/3584931.3607492>
- [88] Yongliang Shen, Kaitao Song, Xu Tan, Dongsheng Li, Weiming Lu, and Yueting Zhuang. 2023. HuggingGPT: Solving AI Tasks with ChatGPT and its Friends in Hugging Face. arXiv:2303.17580 [cs.CL]
- [89] Grace Shin, Yuanqun Feng, Mohammad Hossein Jarrahi, and Nicci Gafinowitz. 2018. Beyond novelty effect: a mixed-methods exploration into the motivation for long-term activity tracker use. *JAMIA Open* 2, 1 (2018), 62–72. <https://doi.org/10.1093/jamiaopen/ooy048> eprint: <https://academic.oup.com/jamiaopen/article-pdf/2/1/62/32298485/ooy048.pdf>.
- [90] Sungbok Shin, Sanghyun Hong, and Niklas Elmqvist. 2023. Perceptual Pat: A Virtual Human Visual System for Iterative Visualization Design. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. ACM, Hamburg Germany, 1–17. <https://doi.org/10.1145/3544548.3580974>
- [91] Nikhil Singh, Guillermo Bernal, Daria Savchenko, and Elena L. Glassman. 2022. Where to Hide a Stolen Elephant: Leaps in Creative Writing with Multimodal Machine Intelligence. *ACM Trans. Comput.-Hum. Interact.* (feb 2022). <https://doi.org/10.1145/3511599> Just Accepted.
- [92] Kihoon Son, DaEun Choi, Tae Soo Kim, Young-Ho Kim, and Juho Kim. 2023. GenQuery: Supporting Expressive Visual Search with Generative Models. <https://doi.org/10.48550/arXiv.2310.01287> arXiv:2310.01287 [cs]
- [93] Adam J Sporka, Torsten Felzer, Sri H Kurmiawan, Ondřej Poláček, Paul Haiduk, and I Scott MacKenzie. 2011. Chanti: Predictive text entry using non-verbal vocal input. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. 2463–2472.
- [94] Sangho Suh, Meng Chen, Bryan Min, Toby Jia-Jun Li, and Haijun Xia. 2024. Luminate: Structured Generation and Exploration of Design Space with Large Language Models for Human-AI Co-Creation. (2024), 1–26. <https://doi.org/doi/10.1145/3613904.3642400>
- [95] Sangho Suh, Bryan Min, Srishti Palani, and Haijun Xia. 2023. Sensecape: Enabling Multilevel Exploration and Sensemaking with Large Language Models. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology* (UIST). <https://doi.org/doi/10.1145/3586183.3606756>
- [96] Symplur. 2019. Tweetorials — From Early Beginnings to Huge Growth and Beyond. <https://www.symplur.com/blog/tweetorials-from-early-beginnings-to-huge-growth-and-beyond/>
- [97] Tiffany Tseng, Ruijia Cheng, and Jeffrey Nichols. 2024. Keyframer: Empowering Animation Design using Large Language Models. arXiv:2402.06071 [cs] <http://arxiv.org/abs/2402.06071>
- [98] Alexandra Voit, Jasmin Niess, Caroline Eckerth, Maïke Ernst, Henrike Weingärtner, and Paweł W. Woźniak. 2020. 'It's not a romantic relationship': Stories of Adoption and Abandonment of Smart Speakers at Home. In *Proceedings of the 19th International Conference on Mobile and Ubiquitous Multimedia* (Essen, Germany) (MUM '20). Association for Computing Machinery, New York, NY, USA, 71–82. <https://doi.org/10.1145/3428361.3428469>
- [99] Sitong Wang, Samia Menon, Tao Long, Keren Henderson, Dingzeyu Li, Kevin Crowston, Mark Hansen, Jeffrey V. Nickerson, and Lydia B. Chilton. 2024. Reel-Framer: Human-AI Co-Creation for News-to-Video Translation. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*. 1–20.
- [100] Sitong Wang, Savvas Petridis, Taeahn Kwon, Xiaojuan Ma, and Lydia B Chilton. 2023. PopBlends: Strategies for conceptual blending with large language models. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. 1–19.
- [101] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc Le, and Denny Zhou. 2023. Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. arXiv:2201.11903 (2023). <https://doi.org/10.48550/arXiv.2201.11903> [cs]
- [102] Spencer Williams, Ridley Jones, Katharina Reinecke, and Gary Hsieh. 2022. An HCI Research Agenda for Online Science Communication. *Proceedings of the ACM on Human-Computer Interaction* 6, CSCW2 (Nov 2022), 1–22. <https://doi.org/10.1145/3555591>
- [103] Qingyun Wu, Gagan Bansal, Jieyu Zhang, Yiran Wu, Beibin Li, Erkang Zhu, Li Jiang, Xiaoyun Zhang, Shaokun Zhang, Jiale Liu, Ahmed Hassan Awadallah, Ryan W White, Doug Burger, and Chi Wang. 2023. AutoGen: Enabling Next-Gen LLM Applications via Multi-Agent Conversation. arXiv:2308.08155 [cs.AI]
- [104] Tongshuang Wu, Ellen Jiang, Aaron Donsbach, Jeff Gray, Alejandra Molina, Michael Terry, and Carrie J Cai. 2022. PromptChainer: Chaining Large Language Model Prompts through Visual Programming. In *Extended Abstracts of the 2022 CHI Conference on Human Factors in Computing Systems* (New Orleans, LA, USA) (CHI EA '22). Association for Computing Machinery, New York, NY, USA, Article 359, 10 pages. <https://doi.org/10.1145/3491101.3519729>
- [105] Tongshuang Wu, Michael Terry, and Carrie Jun Cai. 2022. AI Chains: Transparent and Controllable Human-AI Interaction by Chaining Large Language Model Prompts. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (New Orleans, LA, USA) (CHI '22). Association for Computing Machinery, New York, NY, USA, Article 385, 22 pages. <https://doi.org/10.1145/3491102.3517582>
- [106] Tongshuang Wu, Haiyi Zhu, Maya Albarak, Alexis Axon, Amanda Bertsch, Wenxing Deng, Ziqi Ding, Bill Guo, Sireesh Gururaj, Tzu-Sheng Kuo, Jenny T. Liang, Ryan Liu, Ihita Mandal, Jeremiah Milbauer, Xiaolin Ni, Namrata Padmanabhan, Subhashini Ramkumar, Alexis Sudjianto, Jordan Taylor, Ying-Jui Tseng, Patricia Vaidos, Zhijin Yu, Wei Wu, and Chenyang Yang. 2023. LLMs as Workers in Human-Computational Algorithms? Replicating Crowdsourcing Pipelines with LLMs. arXiv:2307.10168 [cs] <http://arxiv.org/abs/2307.10168>
- [107] Shishi Xiao, Liangwei Wang, Xiaojuan Ma, and Wei Zeng. 2024. TypeDance: Creating Semantic Typographic Logos from Image through Personalized Generation. <https://doi.org/10.1145/3613904.3642185> arXiv:2401.11094 [cs]
- [108] Yuxi Xie, Kenji Kawaguchi, Yiran Zhao, Xu Zhao, MingSung Kan, Junxian He, and Qizhe Xie. 2023. Self-Evaluation Guided Beam Search for Reasoning. <https://api.semanticscholar.org/CorpusID:258426922>
- [109] Kevin Yang, Yuandong Tian, Nanyun Peng, and Dan Klein. 2022. Re3: Generating Longer Stories With Recursive Reprompting and Revision. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing* (Abu Dhabi, United Arab Emirates, 2022-12), Yoav Goldberg, Zornitsa Kozareva, and Yue Zhang (Eds.). Association for Computational Linguistics, 4393–4479. <https://doi.org/10.18653/v1/2022.emnlp-main.296>
- [110] J.D. Zamfirescu-Pereira, Richmond Y. Wong, Bjoern Hartmann, and Qian Yang. 2023. Why Johnny Can't Prompt: How Non-AI Experts Try (and Fail) to Design LLM Prompts. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (Hamburg, Germany) (CHI '23). Association for Computing Machinery, New York, NY, USA, Article 437, 21 pages. <https://doi.org/10.1145/3544548.3581388>
- [111] Qianqia (Queenie) Zhang, Soya Park, Michael Muller, and David R. Karger. 2024. "How fancy you are to make us use your fancy tool": Coordinating Individuals' Tool Preference over Group Boundaries. *Proc. ACM Hum.-Comput. Interact.* 8, GROUP, Article 4 (feb 2024), 31 pages. <https://doi.org/10.1145/3633069>
- [112] Zheng Zhang, Jie Gao, Ranjodh Singh Dhaliwal, and Toby Jia-Jun Li. 2023. VISAR: A Human-AI Argumentative Writing Assistant with Visual Programming and Rapid Draft Prototyping. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology* (San Francisco, CA, USA) (UIST '23). Association for Computing Machinery, New York, NY, USA, Article 5, 30 pages. <https://doi.org/10.1145/3586183.3606800>
- [113] Ran Zhou, Zachary Schwemler, Akshay Baweja, Harpreet Sareen, Casey Lee Hunt, and Daniel Leithinger. 2023. TactorBots: A Haptic Design Toolkit for Out-of-lab Exploration of Emotional Robotic Touch. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. ACM, Hamburg Germany, 1–19. <https://doi.org/10.1145/3544548.3580799>
- [114] Zhongyi Zhou, Jing Jin, Vrushank Phadnis, Xiuxian Yuan, Jun Jiang, Xun Qian, Jingtao Zhou, Yiyi Huang, Zheng Xu, Yinda Zhang, Kristen Wright, Jason Mayes, Mark Sherwood, Johnny Lee, Alex Olwal, David Kim, Ram Iyengar, Na Li, and Ruofei Du. 2023. InstructPipe: Building Visual Programming Pipelines with Human Instructions. arXiv:2312.09672 [cs] <http://arxiv.org/abs/2312.09672>

A COMPLETE SYSTEM WALKTHROUGH

Step 1: Input Topic & Generate Everyday Examples

[tweet_topic]:

language models

prompt:

Please list 5 everyday examples of [tweet_topic]. Each example should be under 10 words, related to a general audience, and not require technical knowledge.

START

Here is a Quick Hook Generation:

Ever wonder how your phone predicts what you're about to type next? 🤖 That's thanks to language models! Think of it as your device's mini fortune-teller, guessing your next word based on what you've typed so far. But how can it predict your personal vocab? Look out for our upcoming blog to dive into the mystery! 🕵️ #LanguageModels #AI #LetsTalkTech

Here are 5 everyday examples:

- Auto-correcting typos on your smartphone. ❤️
- Predictive texting while writing messages. ❤️ ↓ Copied to Next Step
- Voice recognition like Siri or Alexa. ❤️
- Search engine suggestions. ❤️
- Email spam filters. ❤️

REGENERATE

Step 2: Generate Common Experiences

[everyday_example]:

Predictive texts and plot suggestions for creative writing

prompt:

Please list 5 common experiences with using [everyday_example]. Each experience should be under 10 words, related to a general audience, and not require technical knowledge.

GENERATE COMMON EXPERIENCES

Here are 5 common experiences:

- Unexpectedly humorous or nonsensical suggestions. ❤️ ↓ Copied to Next Step
- Autocorrect mistakes changing intended meaning. ❤️

- Helpful in faster typing and spelling correction. ❤️
- Stunting creative expression with standard phrases. ❤️
- Learning to use unusual words or phrases less often. ❤️

REGENERATE

Step 3: Generate User Scenarios

[common_experience]:

Unexpectedly humorous or nonsensical suggestions.

prompt:

Please list 12 user scenarios someone encounters with [common_experience] while using [tweet_topic]. Each scenario should be under 15 words, from a first-person perspective, related to a general audience, and not require technical knowledge.

GENERATE USER SCENARIOS

Here are 5 user scenarios:

- While writing an email, it suggested I end with "Roar like a dinosaur." ❤️
- It replaced "Good morning" with "Party with unicorns." ❤️
- It suggested to rephrase "Will you marry me?" with "Can you be my penguin?" ❤️
- Using it as a recipe guide, it added "dance salsa for 5 mins" in cooking steps. ❤️
- It replaced "meeting at 3pm" with "mammoth cave tour at sunset." ❤️
- When typing "Guardians of", it completed as "Guardians of the gummy bears." ❤️ ↓ Copied to Next Step
- While typing a complaint letter, it suggested "Let's have a pizza party." ❤️
- Asked about the weather, it replied me with "Cloudy chance of flying saucers." ❤️
- Documenting meeting minutes, it recommended "Meeting was an alien conspiracy." ❤️
- It suggested citing "Report released by Invisible College of Fairy Engineers" in my assignment. ❤️
- Writing a heartfelt poem, it ended with "And so, the broccoli whispered." ❤️
- Asked for the synonym of "happy", it suggested "cheerful cheerleaders" ❤️

REGENERATE

Step 4: Generate Personal Anecdotes

[user_scenario]:

When typing "Guardians of", it completed as "Guardians of the gummy bears"

prompt:

Please write 3 short personal anecdotes about the topic, [tweet_topic], for the user scenario of [user_scenario] for creative writing. Each anecdote should be under 30 words, from a creative writer's perspective, related to a general audience, and not require technical knowledge. Make sure to mention that the writer actually incorporates the funny suggestion for their future plot.

GENERATE PERSONAL ANECDOTES

Here are the 3 personal anecdotes:

1. While crafting a superhero novel, I typed "Guardians of..." intending to type "Guardians of the Galaxy", but my language model amusingly suggested "Guardians of the Gummy Bears". It sparked an idea for a humorous subplot.

2. I was creating a mythical saga and paused at "Guardians of...". Hilariously, the language model completed it as "Guardians of the Gummy Bears". I worked it in as a fantastical microcosm.

3. Absentmindedly, pasting "Guardians of..." into a political satire, the language model proposed "Guardians of the Gummy Bears". I laughed out loud and turned it into a whimsical political party.
- REGENERATE

Step 5: Generate More Specific Anecdotes

[personal_anecdote]:

While crafting a superhero novel, I typed "Guardians of..." intending to type "Guardians of the Galaxy", but my language model amusingly suggested "Guardians of the Gummy Bears". It sparked an idea for a humorous subplot.

prompt:

Please rewrite 3 versions of this anecdote to make it more specific. Replace all generic terms with concrete and specific instances: [personal_anecdote].

GENERATE SPECIFIC ANECDOTES

Here are 3 specific versions:

1. While I was diligently typing up my latest Batman-inspired book, my fingers swiftly moved to type the familiar phrase "Guardians of Gotham City". However, my cheeky automated spellcheck unexpectedly suggested "Guardians of the Gumball Machines". This surprise sparked an intriguing concept for a side-splitting sideline to my main plot.

2. Once, during an intense session of composing my Marvel-based masterpiece, I began to write "Guardians of the Cosmic Gems", calling back to my cherished "Guardians of the Galaxy" characters. But my computer, with its AI language model, hilariously replaced it with "Guardians of the Candy Canes". This hilarious mix-up sparked a brilliant idea for an uproariously funny tangent within my main story.

3. One day, while I was passionately crafting an Avengers-centric narrative, I was all set to reference the iconic "Guardians of the Infinity Stones". Suddenly, my quirky autocorrect amusingly provided the suggestion "Guardians of the Jelly Beans". This humorous curveball ignited an idea for a comical thread weaving its way through the main narrative.
- REGENERATE

Step 6: Generate Example Hooks

[specific_anecdote]:

Last week, while crafting a sci-fi novel, I typed "Guardians of..." intending to type "Guardians of the Galaxy", but my language model amusingly suggested "Guardians of the Gummy Bears". It sparked an idea for a humorous subplot. I mused: WHY NOT MAKE IT SWEET? And with a giggle, I wrote a whole chapter about the gummy bear kingdom. 🍬🐻👑

prompt:

Here is a personal anecdote about the topic, [tweet_topic], for the experience of [common_experience]: [specific_anecdote].

Based on this anecdote, please write 3 versions of the opening tweet to introduce the concept of [tweet_topic]. The tweet should have no jargon. It should contain a specific and relatable example. It should have a driving question to spark curiosity and make readers want to read more. Make sure to include a few emojis and all the details.

GENERATE EXAMPLE HOOKS!

Here are 3 example hooks generated by GPT:

1. Ever typed "Guardians of" and expected "the Galaxy" but got "the Gummy Bears" instead? 🤖 What if I told you, this was no random suggestion but the workings of a brilliant technology called "Language Models"? Got your interest piqued? Then follow me into this enlightening thread! 🧵1/

2. As I penned down an epic sci-fi novel last week, my writing assistant threw a curveball by suggesting "Guardians of the Gummy Bears" 🐻🍬 Can you believe it? It made me chuckle, as well as ponder. So, what's generating these unexpected yet fun suggestions? Venture with me into the realm of "Language Models". / 1/

3. Picture this: You're mindlessly keyboarding away, and out pops a suggestion so silly, it makes your day! Like when I was drafting my sci-fi saga and my device suggests "Guardians of the Gummy Bears"! 🐻🍬 Intrigued by these funny suggestions? Let's unravel the magic of "Language Models" together. 1/

REGENERATE

Step 7: Write a Final Hook

Ever typed "Guardians of" and expected "the Galaxy" but got "the Gummy Bears" suggested? 🤖🐻🍬 I just got sugary space heroes featured in my new sci-fi! SO CUTE! 🧴🧴! What if I told you, this was no random suggestion but a brilliant technology called "Language Models"? 🧵1/

275/280

TWEET

B WEEK 3 SYSTEM UPDATES: STEP 6 LEARN-MY-STYLE PROMPT

Step 6

Week 1 - 2 (session 1 - 7)

Step 6: Generate Example Hooks

[specific_anecdote]:

Please enter the specific personal anecdote here.

prompt:

[original_training_examples_dont_delete_me]

Here is a personal anecdote about the topic, [tweet_topic], for the experience of [common_experience]: [specific_anecdote].

Based on this anecdote, please write 3 versions of the opening tweet to introduce the concept of [tweet_topic]. The tweet should have no jargon. It should contain a specific and relatable example. It should have a driving question to spark curiosity and make readers want to read more.

GENERATE EXAMPLE HOOKS!

Step 6 with the Learn-My-Style prompt

Week 3: (session 8 - 10)

Step 6: Generate Example Hooks

[specific_anecdote]:

Please enter the specific personal anecdote here.

prompt:

Original Prompt Learn-My-Style Prompt

[original_training_examples_dont_delete_me]

Here is a personal anecdote about the topic, [tweet_topic], for the experience of [common_experience]: [specific_anecdote].

Based on this anecdote, please write 3 versions of the opening tweet to introduce the concept of [tweet_topic]. The tweet should have no jargon. It should contain a specific and relatable example. It should have a driving question to spark curiosity and make readers want to read more.

GENERATE EXAMPLE HOOKS!

Step 6: Generate Example Hooks

[specific_anecdote]:

Please enter the specific personal anecdote here.

prompt:

Original Prompt Learn-My-Style Prompt

Here are examples of interesting opening tweets to introduce various computer science topics:

Operating System: 🧠 If your computer's got just one brain (one CPU), how's it simultaneously emailing, playing tunes, and downloading movies? 🤖 Unveil the mystery behind this digital circus! Join me in unraveling the captivating realm of Operating Systems and their awe-inspiring juggling prowess. 🎪

Memory Disaggregation! Ever had your phone crash while editing a video because it runs out of memory, meanwhile your tablet is lying idle with ample memory? 📱💻 What if your phone could borrow some of that? 🤖 What's stopping them? 🤖 Let me introduce you to memory disaggregation, a concept which talks about this... 1/ 📄

Here is a personal anecdote about the topic, [tweet_topic], for the experience of [common_experience]: [specific_anecdote].

Based on this anecdote, please write 3 versions of the opening tweet to introduce the concept of [tweet_topic]. The tweet should have no jargon. It should contain a specific and relatable example. It should have a driving question to spark curiosity and make readers want to read more.

GENERATE EXAMPLE HOOKS!