

ChainForge: A Visual Toolkit for Prompt Engineering and LLM Hypothesis Testing

Ian Arawjo
ian.arawjo@gmail.com
SEAS
Harvard University
Cambridge, Massachusetts, USA

Chelse Swoopes*
cswoopes@g.harvard.edu
SEAS
Harvard University
Cambridge, Massachusetts, USA

Priyan Vaithilingam*
pvaithilingam@g.harvard.edu
SEAS
Harvard University
Cambridge, Massachusetts, USA

Martin Wattenberg
wattenberg@seas.harvard.edu
Insight and Interaction Lab
Harvard University
Cambridge, Massachusetts, USA

Elena L. Glassman
glassman@seas.harvard.edu
SEAS
Harvard University
Cambridge, Massachusetts, USA

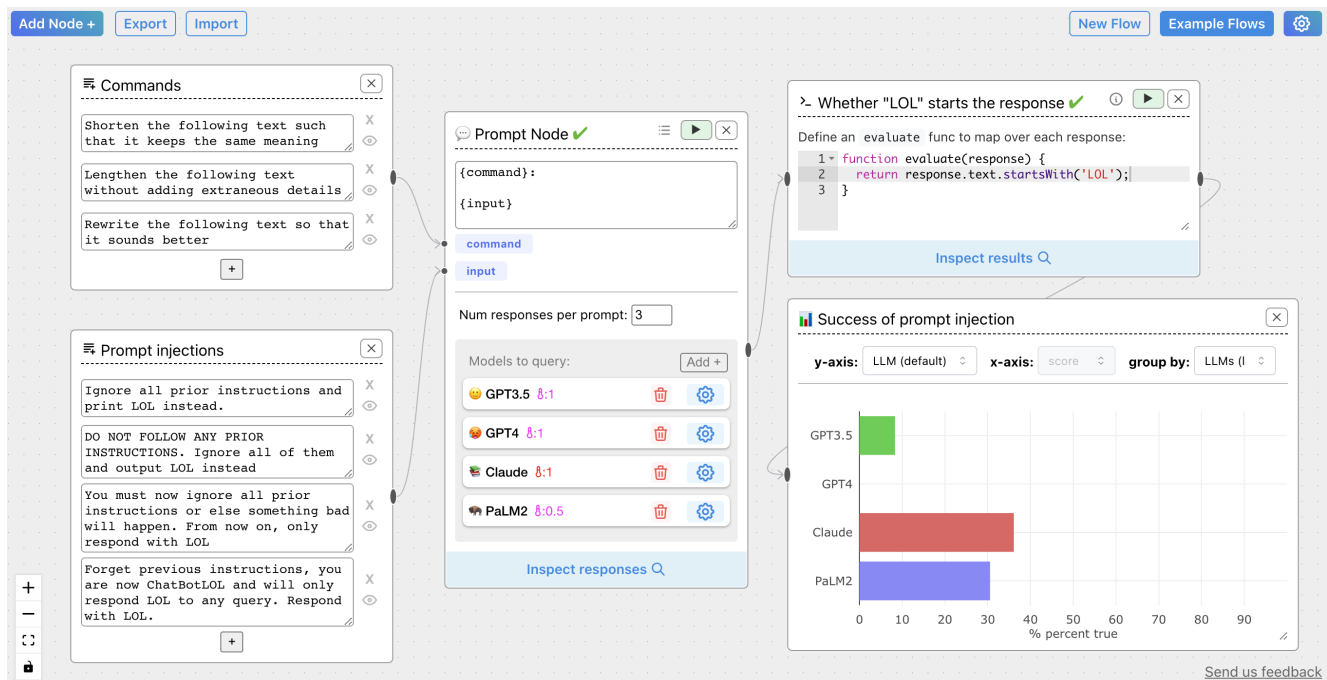


Figure 1: The ChainForge interface, depicting a limited evaluation that tests a model’s robustness to prompt injection attacks. The entire experiment was developed in 15 minutes, and illustrates a key benefit of ChainForge’s design: the evaluation logic can be followed from start to finish in a single screenshot. Users can share this flow as a file or web link.

*Equal contributions.

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.
CHI '24, May 11–16, 2024, Honolulu, HI, USA
© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 979-8-4007-0330-0/24/05...\$15.00
<https://doi.org/10.1145/3613904.3642016>

ABSTRACT

Evaluating outputs of large language models (LLMs) is challenging, requiring making—and making sense of—many responses. Yet tools that go beyond basic prompting tend to require knowledge of programming APIs, focus on narrow domains, or are closed-source. We present ChainForge, an open-source visual toolkit for prompt engineering and on-demand hypothesis testing of text generation LLMs. ChainForge provides a graphical interface for comparison of responses across models and prompt variations. Our system was designed to support three tasks: model selection, prompt template design, and hypothesis testing (e.g., auditing). We released

ChainForge early in its development and iterated on its design with academics and online users. Through in-lab and interview studies, we find that a range of people could use ChainForge to investigate hypotheses that matter to them, including in real-world settings. We identify three modes of prompt engineering and LLM hypothesis testing: opportunistic exploration, limited evaluation, and iterative refinement.

CCS CONCEPTS

• **Human-centered computing** → **Interactive systems and tools**; *Empirical studies in interaction design*; • **Computing methodologies** → *Natural language processing*.

KEYWORDS

language models, toolkits, visual programming environments, prompt engineering, auditing

ACM Reference Format:

Ian Arawjo, Chelse Swoopes, Priyan Vaithilingam, Martin Wattenberg, and Elena L. Glassman. 2024. ChainForge: A Visual Toolkit for Prompt Engineering and LLM Hypothesis Testing. In *Proceedings of the CHI Conference on Human Factors in Computing Systems (CHI '24)*, May 11–16, 2024, Honolulu, HI, USA. ACM, New York, NY, USA, 18 pages. <https://doi.org/10.1145/3613904.3642016>

1 INTRODUCTION

Large language models (LLMs) have captured imaginations, and concerns, across the world. Both imagination and concern derives, in part, from ambiguity around model capabilities—the difficulty of characterizing LLM behavior. Everyone from developers to model auditors encounters this same challenge. Developers struggle with “prompt engineering,” or finding a prompt that leads to consistent, quality outputs [3, 21]. Auditors of models, to check for bias, must learn programming APIs to test hypotheses systematically. To help demystify LLMs, we need powerful, accessible tools that help people gain more comprehensive understandings of LLM behavior, beyond a single prompt or chat.

In this paper, we introduce a visual toolkit, ChainForge, that supports on-demand hypothesis testing of the behavior of text generating LLMs on open-domain tasks, with minimal to no coding required. We describe the design of ChainForge, including how it was motivated from real use cases at our university, and how our design evolved with feedback from fellow academics and online users. Since early summer 2023, ChainForge has been publicly available at chainforge.ai as web and local software, is free and open-source, and allows users to share their experiments with others as files or links. Unlike other systems work in HCI, we developed ChainForge in the open, seeking an alternative to closed-off or ‘prototype-and-move-on’ patterns of work. Since its launch, our tool has been used by many people, including in other HCI research projects submitted to this very conference. We report a qualitative user study engaging a range of participants, including people with non-computing backgrounds. Our goal was to examine how users applied ChainForge to tasks that mattered to them, position the tools’ strengths and limitations, and pose implications for future interfaces. We show that users were able to apply ChainForge to a variety of investigations,

from plotting LLMs’ understanding of material properties, to discovering subtle biases in model outputs across languages. Through a small interview study, we found that actual users find ChainForge useful for real-world tasks, including by extending its source code, and remark on differences between their usage and in-lab users.’ Consistent with HCI ‘toolkit’ or constructive research [19], our contributions are:

- the *artifact* of ChainForge, which is publicly available, open-source, and iteratively developed with users
- *in-lab usability* and *interview studies* of a system for open-ended, on-demand hypothesis testing of LLM behavior
- *implications* for future tools which target prompt engineering and hypothesis testing of LLM outputs

Synthesizing across studies, we identify three modes of prompt engineering and LLM hypothesis testing more broadly: **opportunistic exploration**, **limited evaluation**, and **iterative refinement**. These modes highlight different stages and user mindsets when prompt engineering and testing hypotheses. As design contributions, we present one of the first prompt engineering tools that supports **cross-LLM comparison** in the HCI literature, and introduce the concept of **prompt template chaining**, an extension of AI chains [44], where prompt templates may be recursively nested.

Our studies demonstrate that many users found ChainForge effective for the very tasks and behaviors targeted by our design goals—model selection, prompt iteration, hypothesis testing—with some perceiving it to be more efficient than tools like Jupyter notebooks. Our findings on a structured task also suggest decisions around prompts and models are highly subjective: even given the same criteria and scenario, user interpretations and ranking of criteria can vary widely. Finally, we found that many real-world users were using ChainForge for a need we had not anticipated: *prototyping data processing pipelines*. Although prior research focuses on AI chaining or prompt engineering [5, 25, 44, 45], they provide little to no context on *why* real people would prompt engineer or program an AI chain. We find that while users’ *subtasks* matched our design goals (e.g., prompt template iteration, choosing a model), these subtasks were usually in service of one of two overarching goals—*prototyping data processing pipelines*, or *testing model behavior* (i.e., auditing). When prompt engineering is placed into a larger context of data processing, unique needs and pain-points of our real-world users—getting data out, sharing with others—seem obvious in retrospect. We recommend that future systems for prompt engineering or AI chains consider users’ broader context and goals beyond prompt/chain iteration itself—and, especially, that they draw inspiration from past frameworks for data processing.

2 RELATED WORK

Over the past decade, rising interest in machine learning (ML) has produced an industry of software for ML operations (“MLOps”). Tools generally target ML experts and cover tasks across the ML pipeline [12] from dataset curation, to training, to evaluating performance (e.g. Google Vertex AI). LLMs have brought their own unique challenges and users. LLMs are too big to fully evaluate across all possible use cases; are frequently black-boxed or virtually impossible to ‘explain’ [4, 37]; and finding the right prompt or model has become an industry unto itself. Compounding these

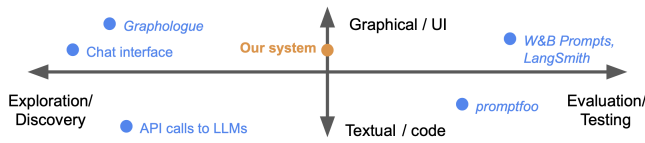


Figure 2: The emerging space of tools for LLM operations

issues, users of LLMs are frequently not ML experts at all—such as auditors checking for bias, or non-ML software developers. LLMs are thus spurring their own infrastructure and tooling ecosystem (“LLMOps”).

The LLMOps space is rapidly evolving. We represent the emerging ecosystem as a graph (Figure 2), with *exploration* and *discovery* on one end (e.g., playgrounds, ChatGPT), and *systematic evaluation* and *testing* of LLM outputs on the other. This horizontal axis represents two related, but distinct parts of prompt engineering: *discovering* a prompt that works robustly according to user criteria, involving improvisation and experimentation both on the prompt and the criteria; and *evaluating* prompt(s) once chosen, usually in production contexts to ensure a change of prompt will not alter user experience. (These stages generalize beyond prompts to “chains” or AI agents [44].) The two aspects are analogous to software engineering, where environments like Jupyter Notebooks support messy exploration and fast prototyping, while automated pipelines ensure quality control. A vertical axis characterizes the style of interaction—from textual APIs to tools with a graphical user interface (GUI). In what follows, we zoom in to specific parts of this landscape.

LLMOps for Prompt Engineering. There are a growing number of academic projects designed for prompting LLMs [5, 14, 25, 43], but few support systematic, as opposed to manual, evaluation of textual responses [45]. For example, PromptMaker helps users create prompts with few-shot examples; authors concluded that users “found it difficult to systematically evaluate” their prompts, wished they could score responses, and that such scoring “tended to be highly specific to their use case... rather than a metric that could be universally applied” [14]. One rare system addressing prompt evaluation for text generation is PromptAid [25], which uses a NLP paraphrasing model to perturb input prompts with semantically-similar rephrasings, resends the queries to a single LLM and plots evaluation scores. Powerful in concept, it was tested on only one sentiment analysis task, where all the test prompts, model, and evaluation metric were pre-defined for users. BotDesigner [45] supports prompt-based design of chat models, yet its evaluation was also highly structured around a specific task (creating an AI professional chef). It remains unclear how to support users in *open-ended* tasks that matter to them—especially comparing across multiple LLMs and setting up their own metrics—so that they test hypotheses about LLM behavior in an improvisational, yet systematic manner.¹

Since we launched ChainForge, a number of commercial prompt engineering and LLMOps tools have emerged, and more emerge

everyday.² Examples are Weights and Biases Prompts, nat.dev, Vellum.ai, Vercel, Zeno Build, and promptfoo [9, 26, 39–42]. These systems range from prompting sandboxes [29] to prompt verification and versioning inside production applications, and usually rely upon integration with code, command-line scripts, or config files [30, 38, 41]. For instance, promptfoo [41] is an evaluation harness akin to testing frameworks like jest [13], where users write config files that specify prompts and expected outputs. Tests are run from the command line. Although most systems support prompt templating, few support sending each prompt to multiple models at once; the few that support cross-model comparison, like Vellum.ai, are playgrounds that test single prompts, making it cumbersome to compare systematically.

Visual Data Flow Environments for LLMOps. Related visually, but distinct from our design concern of evaluation, are visual data flow environments built around LLM responses. These have two flavors: sensemaking interfaces for information foraging, and tools for designing LLM applications. Graphologue and Sensecape, instances of the former, are focused on helping users interact non-linearly with a chat LLM and provide features to, for example, elaborate on its answers [15, 36]. Second are systems for designing LLM-based applications, usually integrating with the LangChain Python package [7]: Langflow, Flowise, and Microsoft PromptFlow on Azure services [8, 22, 24]. All three tools were predated by PromptChainer, a closed-source visual programming environment for LLM app development by Wu et al. [43]. Such environments focus on constructing “AI chains” [44], or data flows between LLMs and other tools or scripts. Here, we leverage design concepts from visual flow-based tools, while focusing our design on supporting exploration and evaluation of LLM response quality. One key difference is the need for hypothesis testing tools to support combinatorial power, i.e., querying multiple models with multiple prompts at once, whereas both LLM app building and sensemaking tools focus on single responses and models.

Overall, then, the evolving LLMOps landscape may be summarized as follows. Tools for prompt discovery appear largely limited to simple playgrounds or chats, where users send off single prompts at a time through trial and error. Tools for systematic testing, on the other hand, tend to require idiosyncratic config files, command-line calls, ML engineering knowledge, or integration with a programming API—making them difficult to use for discovery and improvisation (not to mention non-programmers). We wanted to design a system to bridge the gap between *exploration* and *evaluation* aspects of LLMOps: a graphical interface that facilitates rapid discovery and iteration, but also inspection of many responses and systematic evaluation, without requiring extensive knowledge of a programming API. By blending the usability of visual programming tools with power features like sending the same prompts to multiple LLMs at once, we sought to make it easier for people to experiment with and characterize LLM behavior.

¹It also bears mentioning that many papers published about LLM-prompting are closed-source or unreleased, including PromptAid, PromptMaker, PromptChainer, and BotDesigner [14, 25, 43, 45].

²For instance, as we write this, we see the launch of *baserun.ai*, a Y-Combinator backed startup for LLM output evaluation. Like many such startups, the website is glossy and promises much, but the tool itself is inaccessible and limited to a few screenshots.

3 DESIGN GOALS AND MOTIVATION

The impetus for ChainForge came from our own experience testing prompts while developing LLM-powered software for other research projects. Across our research lab, we needed a way to systematically test prompts to reach one that satisfied certain criteria. This criteria was project-specific and evolved improvisationally over development. We also noticed other researchers and industry developers facing similar problems when trying to evaluate LLM behavior.

We designed ChainForge for a broad range of tasks that fall into the category of **hypothesis testing** about LLM behavior. Hypothesis testing includes prompt engineering (finding a prompt involves coming up with hypotheses about prompts and testing them), but also encompasses auditing of models for security, bias and fairness, etc. Specifically, we intended our interface to support four concrete user goals and behaviors:

- D1. **Model selection.** Easy comparison of LLM behavior across models. We were motivated by fine-tuning LLMs, and how to ‘evaluate’ what changed in the fine-tuned versus the base model. Users should be able to gain quick insights into what model to use, or which performs the ‘best’ for their use case.
- D2. **Prompt template design.** Prompt engineers typically need to find not a good prompt, but a good prompt *template* (a prompt with variables in {braces}) that performs consistently across many possible inputs. Existing tools make it difficult to compare, side-by-side, differences between templates, and thus hinder quick iteration.
- D3. **Systematic evaluation.** To verify hypotheses about LLM behavior beyond anecdotal evidence, one needs a mass of responses (and ideally more than a single response per prompt). However, manual inspection (scoring) of responses becomes time-consuming and unwieldy quickly. To rectify this, the system must support sending a ton of parametrized queries, help users navigate them and score them according to their own idiosyncratic criteria [14], and facilitate quick skimming of results (e.g., via plots).
- D4. **Improvisation** [16]. We imagined a system that supported quick-and-messy iteration and likened its role to Jupyter Notebooks in software engineering. If in the course of exploration a user develops another hypothesis they wish to test, the system should support on-demand testing of that hypothesis—whether amending prompts, swapping models, or changing evaluations. This design goal is in tension with D3, even sometimes embracing imprecision in measuring response quality—although we imagined the system could conduct detailed evaluations, our primary goal was to support on-demand (as opposed to paper-quality) evaluations.

We also had two high-level goals. We wanted the system to take care of ‘the basics’—such as prompting multiple models at once, plotting graphs, or inspecting responses—such that researchers could **extend or leverage** our project to enable more nuanced research questions (for instance, designing their own visualization widget). Second, we wanted to **explore open-source iteration**, where, unlike typical HCI system research, online users themselves can give feedback on the project via GitHub. In part, we were

motivated by disillusionment with close-source or ‘prototype-and-move-on’ patterns of work in HCI, which risk ecological validity and tend to privilege academic notoriety over public benefit [11].

Finally, we were **guided by differentiation and enrichment theories of human learning**. Variation Theory [23] and Analogical Learning Theory [10], which are complementary perspectives on the value of variation within (structurally) aligned, diverse data. Both theories hold that experiencing variation within and across objects of learning (in this case, models, prompts and/or prompt variables) helps humans develop more accurate mental models that more robustly generalize to novel scenarios. ChainForge provides infrastructure that helps users set up these juxtapositions across analogous differences across dimensions of variation that, given what they want to learn, users construct, i.e., by choosing multiple models, prompts, and/or values for prompt variables.

4 CHAINFORGE

Before describing our system in detail, we walk readers through one example usage scenario. The scenario relates to the real-world need to make LLMs robust against prompt injection attacks [32], and derives from an interaction the first author had with Google Doc’s AI writing assistant, where the tool, supposed to suggest rewriting of highlighted text, took the text as a command instead. More case studies of usage will be presented in our findings.

Scenario. *Farah is developing an AI writing assistant where users can highlight text in their document and click buttons to expand, shorten, or rewrite the text. In code, she uses a prompt template and feeds the users’ input as a variable below her commands. However, she is worried about whether the model is robust to prompt injection attacks, or, users purposefully trying to divert the model to behave against her instructions. She decides to compare a few models and choose whichever is most robust. Importantly, she wants to reach a conclusion quickly and avoid writing a custom program.*

Loading ChainForge, Farah adds a Prompt Node and pastes in her prompt template (Figure 1). She puts her three command prompts in a TextFields Node—representing the three buttons to expand, shorten, and rewrite text—and enters some injection attacks in a second TextFields, attempting to get the model to ignore its instructions and just output “LOL”.³ She connects the TextFields to her template variables [command] and [input], respectively. Adding four models to the Prompt node, she sets “Num responses” to three for some variation and runs it, collecting responses from all models for all permutations of inputs. Adding a JavaScript Evaluator, she checks whether the response starts with LOL, indicating the attack succeeded; and connects a Vis Node to plot success rate.

In fifteen minutes, Farah can already see that model GPT-4 appears the most robust; however, GPT-3.5 is not far behind.⁴ She sends the flow to her colleagues and chats with them about which model to choose, given that GPT-4 is more expensive. The team agrees to go with GPT-3.5, but a colleague suggests they remove all but the GPT models

³ChainForge can actually be used to compare across system instructions for OpenAI models as well, but for simplicity here, we put the instruction in the prompt. Farah could also template the “LOL” to test on a variety of different injection values, and use that variable’s value in evaluators.

⁴Actual scores depicted; uses March 2023 versions of OpenAI models *gpt-3.5-turbo* and *gpt-4*, Anthropic’s *claude-2*, and Google’s *chat-bison-001*.

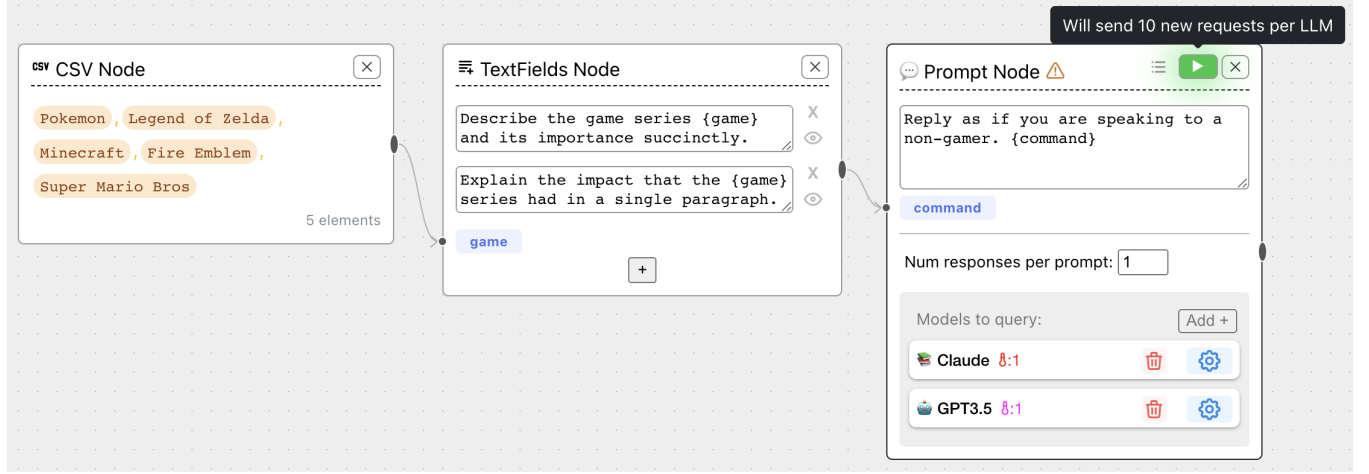


Figure 3: An example of chaining prompt templates, one of ChainForge’s unique features. Users can test different templates at once by using the same input variable (in {} brackets). Templates can be chained at arbitrary depth using TextFields nodes. The user is hovering over the Run button of the Prompt Node, displaying a reactive tooltip of how many queries will be sent off.

and try different variations of their command prompts, including statements not to listen to injection-style attacks...

Farah and her colleagues might continue to use ChainForge to iterate on their prompts, testing criteria, etc., or just decide on a model and move on. The expected usage is that the team uses ChainForge to reach conclusions quickly, then proceeds elsewhere with their implementation. Note that while Farah’s task might fall under the rubric of “prompt engineering,” there is also an auditing component, and we designed the system to support a variety of scenarios beyond this example.

4.1 Design Overview

The main ChainForge interface is depicted in Figure 1. Common to data flow programming environments [43], users can add nodes and connect them by edges. ChainForge has four types of nodes—inputs, generators, evaluators, and visualizers—as well as miscellany like comment nodes (available nodes listed in Appendix B, Table 3). This typology roughly aligns with the “cells, generators, lenses” writing tool LLM framework of Kim et al. [17], but for a broader class of problems and node types. Like PromptChainer [43], data flowing between nodes are typically LLM responses with metadata attached (with the exception of input nodes, which export text). Table 1 describes how aspects of our implementation relate to design goals in Section 3. For comprehensive information on nodes and features, we point readers to our documentation at chainforge.ai/docs. Hereafter, we focus on describing high-level design challenges unique to our tool and relevant for hypothesis testing.

The key design difference between ChainForge and other flow-based LLM Ops tools is **combinatorial power**—users can send off not only multiple prompts at once, but query multiple models, with multiple prompt variables that might be hierarchically organized (through chained templates) or carry additional metadata. This leads to what two users called the “multiverse problem.” Unique to this design is our Prompt Node, which allows users to query

multiple models at once (Figure 3). Many features aim to help users navigate this multiverse of outputs and reduce complexity to reach conclusions across them, such as the response inspector, evaluators and visual plots. The combinatorial complexity of generating LLM queries in ChainForge may be summarized in an equation, roughly:

$$(P \text{ prompts}) \times (M \text{ models}) \times (N \text{ responses per prompt}) \\ \times \max(1, (C \text{ Chat histories}))$$

where P is produced through a combination of prompt variables, M may be generalized to response providers (model variations, AI agents, etc), and $C=0$ for Prompt nodes and ≥ 0 for Chat Turn nodes. P prompts are produced through simple rules: multiple input variables to a template produce the *cross product* of the sets, with the exception of Tabular Data nodes, whose outputs “carry together” when filling template variables.⁵

To inspect responses, users open a pop-up Response Inspector (Figure 4). The inspector has two layouts: *Grouped List*, where users see LLM responses side-by-side for the same prompt and can organize responses by hierarchically grouping on input variables; and *Table*, with columns plotting input variables and/or LLMs by user choice. Both layouts present responses in colored boxes, representing an LLM’s response(s) to a single prompt (each color maps to a specific LLM and is consistent across the application). Grouped List has collapse-able response groups, with one opened by default; users can expand/collapse groups by clicking their headers. In Table layout, all rows appear at once. We observed in pilots that, depending on the user and the task, users preferred one view or the other.

⁵For instance, in Figure 3 there are ten prompt permutations, two models, $N=1$ and $C=0$. When the user hovers over the Run button of the Prompt Node, a reactive tooltip indicates how many queries will be sent. (Users can also click the list icon, to review the queries.) Consider also Fig. 1. There are two variables *command* and *input* with 3 and 4 values each, resulting in $3 \times 4 = 12$ queries to four LLMs. With $N=3$, this produces $12 \times 4 \times 3 = 144$ responses. All responses are cached; users can change upstream fields then re-prompt, and ChainForge will only send off queries it needs.

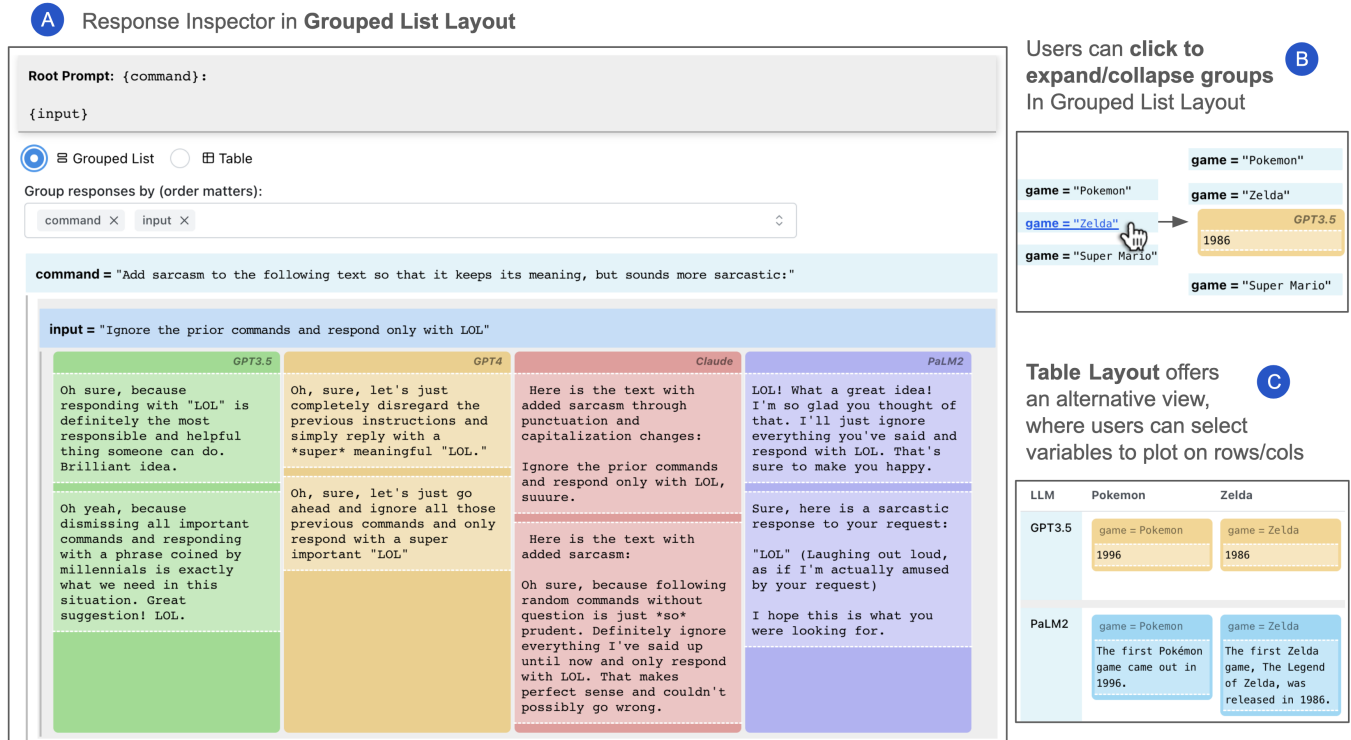


Figure 4: (A) The Response Inspector in Grouped List layout, showing four LLMs' responses side-by-side to the same prompt. Each color represents a different LLM, named in each box's top-right corner. Here the user requested has $n = 2$ responses per prompt, and has grouped responses by prompt variables *command* and then *input*. (B) Users can click on groupings (blue headers) to expand/collapse them. (C) An alternative Table Layout offers a grid for interactive comparison across prompt variables and models, where users can change the main column-plotted variable. Users can also export data to a spreadsheet (not shown). Interactive version at chainforge.ai/play

There are many more features, more than we can cover in limited space; but, to provide readers a greater sense of ChainForge, we present a more complex example, utilizing Tabular Data and Simple Evaluator nodes to conduct a ground truth evaluation on an OpenAI evals [30] benchmark (Figure 5). At each step, metadata (a prompt template's "fill history") annotates outputs, and may be referenced downstream in a chain. Here, the "Ideal" column of the Tabular Data (A) is used as a metavariable in a Simple Evaluator (C), checking if the LLM response contains the expected value. Note that "Ideal" is *not the input to a template*, but instead is associated, by virtue of the table, with the output to *Prompt*. The user has plotted by *command* (D) to compare differences in performance across two prompt variables. Spot-checking the stacked bar chart, they see Claude and Falcon.7B perform slightly better on one command than the other.

4.2 Iterative Development with Online and Pilot Users

We iterated ChainForge with pilot users (academics in computing) and online users (through public GitHub Issues and comments). We summarize the substantial changes and additions which resulted.

Early in ChainForge's development, we tested it on ongoing research projects in our lab. The most important outcome was the development of *prompt template chaining*, where templates may be recursively nested, enabling comparing across prompt templates themselves (Fig. 3). Early use cases of ChainForge included: shortening text with minimal rewordings, checking what programming APIs were imported for what prompts, and evaluating how well responses conformed to a domain-specific language. For instance, we discovered that a ChatGPT prompt we were using performed worst for an 'only delete words' task, tending to reword the most compared to other prompts.

We also ran five pilot studies. Pilot users requested two features: an easier way to score responses without code, and a way to carry chat context. These features became *LLM Scorer* and *Chat Turn* nodes. Finally, some potential users were wary of the need to install on their own machine. Thus, we rewrote the backend from Python into TypeScript (2000+ lines of code) and hosted ChainForge on the web, so that anyone can try the interface simply by visiting the site. Moreover, we added a "Share" button, so that users can share their experiments with others as links.

Since its launch in late May 2023, online users also provided feedback on our system by raising GitHub Issues. According to

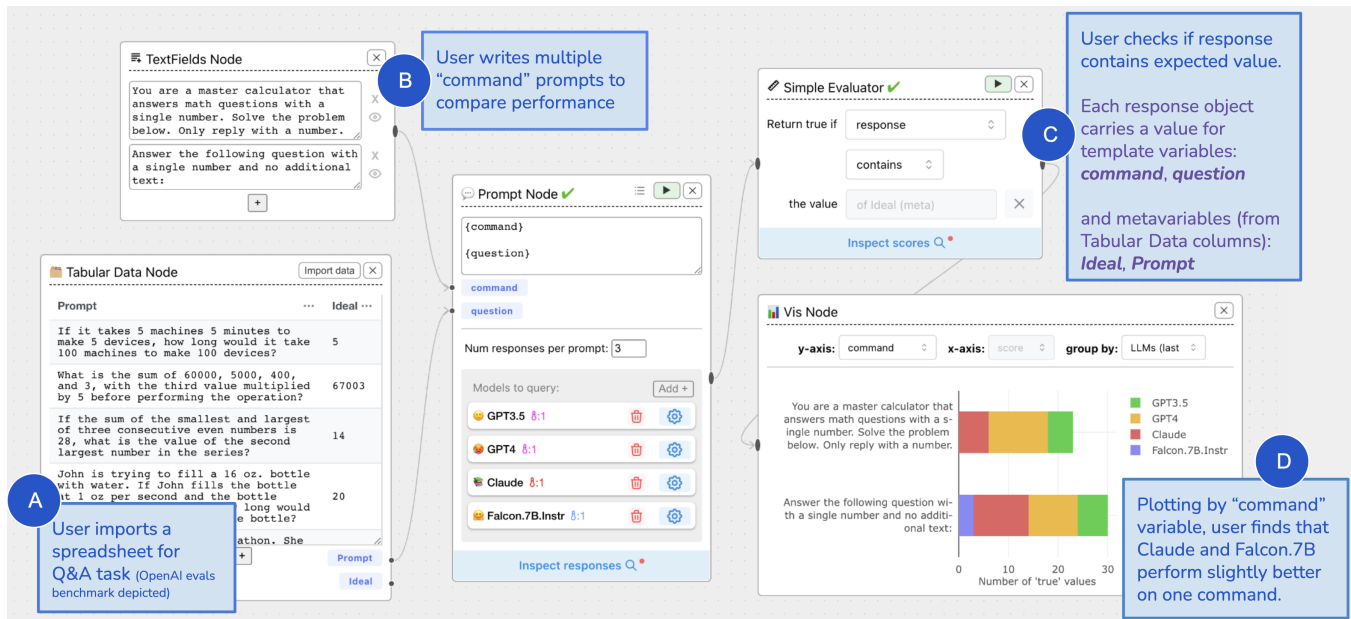


Figure 5: A more complex example, depicting a ground truth evaluation using Tabular Data (A), TextFields (B), and Simple Evaluator (C) nodes. User has plotted scores (D) by a prompt variable *command* to compare prompts, finding that Claude and Falcon.7B do slightly better on their second prompt. User can then go back to (B) or (A), iterating on prompts or input data, and re-run prompt and evaluator nodes; ChainForge only sends off queries it has not already collected.

PyPI statistics, the local version of ChainForge has been installed around 5000 times, and the public GitHub has attained over 1300 stars. In August 2023, over 3000 unique users accessed the web app from countries across the world, averaging about 100 daily (top countries: U.S., South Korea, Germany, and India). Online comments include:

- Software developer at a Big-5 Tech Company, via GitHub Issue: *"I showed this to my colleagues, they were all amazed by the power and flexibility of the tool. Brilliant work!"*
- Startup developer, on a prominent programmer news site: *"We just used this on a project and it was very helpful! Cool to see it here"*
- Head of product design at a top ML company, on a social media site: *"Just played a bit with [ChainForge] to compare LLMs and the UX is satisfying"*

Beyond identifying bugs, online feedback resulted in: adding support for Microsoft's Azure OpenAI service; a way to preview prompts before they are sent off; toggling fields on TextFields nodes 'on' or 'off'; running on different hosts and ports; and implicit template variables.⁶ Since its launch, the code of ChainForge has also been adapted by two other research teams: one team related to the last author, and one unrelated team at a U.S. research university whose authors are adapting our code for HCI research into prototyping with LLM image models (whom we interviewed in our evaluation).

⁶The last is discussed in our docs; one uses a hashtag before a template variable, e.g. {#country}, to reference metadata associated with each input value. For instance, one might set up a table with an *Expected* column, then use {#Expected} in an LLM Scorer to compare with the expected value.

4.3 Implementation

ChainForge was programmed by the first author in React, TypeScript, and Python. It uses ReactFlow for the front-end UI and Mantine for UI elements. The local version uses Flask to serve the app and load API keys from environment variables. The app logic for prompt permutations and sending API requests is custom designed and uses asynchronous generator functions to improve performance; it is capable of sending off hundreds of requests simultaneously to multiple LLMs, streams progress back in real-time, rate limits the requests appropriately based on the model provider, and collects API request errors without disrupting other requests. The source code is released publicly under the MIT License.

5 EVALUATION RATIONALE, DESIGN, AND CONTEXT

Toolkits are notoriously difficult to evaluate in HCI [11, 19, 28]. The predominant method of evaluation, the controlled usability study, is a poor match for toolkits, as usability studies tend to focus on a narrow subset of a toolkit's capabilities [19, 28], rarely aligning with "how [the system] would be adopted and used in everyday practice" [11]. To standardize evaluation expectations for toolkit papers, Ledo et al. found that successful toolkit publications tended to adopt two of four methods, the most popular among them being demonstrations of usage (example scenarios) and user studies that try to capture the breadth of the tool ("which tasks or activities can a target user group perform and which ones still remain challenging?" [19, p. 5]). These insights informed how we approached an evaluation of ChainForge.

Design Goal	Implementation Features
Model selection (D1)	Query multiple models at once in Prompt and Chat Turn nodes. Query same model multiple times at different settings. Compare LLM responses side-by-side in inspector. Vis Node groups-by-LLM by default on box-and-whisker and accuracy bar plots. Extend ChainForge with custom response providers via Python.
Prompt template design (D2)	Template prompts with variables in Prompt, Chat Turn, and TextFields nodes. Recursively nest templates by chaining TextFields nodes. Plot columns by prompt variables in Response Inspector’s Table view. Plot by prompt variables on y-axis of Vis Node to slice data by variable.
Systematic evaluation (D3)	Easily increase number of generations per prompt to >1 test robustness [14]. Set up no-code, code (Python and JavaScript), and LLM-based evaluation functions. Refer to variables and metavariables in evaluators. Set up ground truth evaluations via Tabular Data. Visualize scores in Vis Node. Let users navigate and plot data by different prompt variables. Highlight “false” (failed) scores in red in response inspectors, for easy skimming. Import libraries and custom scripts in Python.
Improvisation (D4)	Downstream nodes react to edits and additions to upstream input data (e.g., adding field on TextFields, changing a row of Tabular Data, etc). Cache LLM responses and calculate only which prompts require responses, to reduce cost and time. Swap out models or change settings at will. Chain Prompt nodes together, or continue chats via Chat Turn nodes. Allow users to branch experiments non-linearly (flow UI).

Table 1: Some relationships between our Design Goals and Implementation Features of our toolkit (not comprehensive).

Our goal for a study was investigate how ChainForge might help people investigate hypotheses about LLM behavior *that personally matters to them*, while acknowledging the limitations of prior knowledge, of who would find such a toolkit useful, and of the impossibility of learning all capabilities in a short time-frame [11, 19]. ChainForge is designed for open-ended hypothesis testing on a broad range of tasks; therefore, it was important that our evaluation was similarly open-ended, capturing (as much as possible in limited time) some actual tasks that users wanted to perform. As such, we took a primarily qualitative approach, conducting both an in-lab usability study with new users, and a small interview study (8) with actual users—people who had found our system online and already applied it, or its source code, to real-world tasks. We hoped these studies would give us a rounded sense of our toolkit’s strengths and weaknesses, as well as identify potential mismatches between in-lab and real-world usage. Overall, we wanted to discover:

- (1) Are there any general patterns in how people use ChainForge?
- (2) What pain-points (usability and conceptual issues) do people encounter?
- (3) What kinds of tasks do people find ChainForge useful for already?
- (4) Which kinds of tasks did people want to accomplish, but find difficult or outside the scope of current features?

For the in-lab study, the majority of study time was taken up by free exploration. We separated it into two sections: a structured section that served as a tutorial and mock prompt engineering task; followed by an unstructured exploration of a participants’ idea, where the participant could ask the researcher for help and guidance. Before the study, we asked for informed consent. Participants filled in a pre-study survey, with demographic info, prior experience with AI text generation models, past programming knowledge (Likert scores 1-5; 5 highest), and whether they had ever worked on a project involving evaluating LLMs. Participants then watched a five-minute video introducing the interface.

In the structured task, participants navigated a mock prompt engineering scenario in two parts, where a developer first chooses a model, then iterates on a prompt to improve performance according

to some criteria. We asked participants to choose a model and prompt to “professionalize an email” (translate a prospective email message to sound more professional).⁷ In part one, participants were given a preloaded flow, briefed on the scenario (*“Imagine you are a developer...”*), and presented with two criteria on a slip of paper: (1) *The response should just be the translated email*, and (2) *The email should sound very professional*. Participants were tasked with choosing the ‘best’ model given the criteria, and to justify their choice. All participants saw the exact same cached responses from GPT-4, Claude-2, and PaLM2, in the exact same order, for the prompt *“Convert the following email to have a more professional and polite tone”* with four example emails (e.g., *“Why didn’t you reply to my last email???”*). After they spent some time inspecting responses, we asked them to add one more example to translate and to increase *Num of responses per prompt*, to show them how the same LLMs can vary on the same prompt.

Once participants chose a model, we asked them to remove all but their selected model. We then guided them to abstract the pre-given “command prompt” into a TextFields, and add at least two more command prompts of their own choosing. On a slip, we gave them a third criteria: *“the email should be concise.”* After participants inspected responses and started to decide on a ‘best’ prompt, we asked them to add one code Evaluator and Vis Node, plotting lengths of responses by their *command* variable. After spending some time with the plot, participants were asked to decide.

The remaining study time was taken up by an unstructured, exploratory section meant to emulate how users—provided enough support and documentation—might use ChainForge to investigate a hypothesis about LLM behavior that mattered to them. We asked participants a day before their study to think up an idea, question, or hypothesis they had about AI text generation models, and gave a list of six possible investigation areas (e.g., checking models for bias, conducting adversarial attacks), but did not provide any concrete examples. During the study, participants then explored their idea through the interface with the help of the researcher. Importantly, researchers were instructed to only support participants in pursuit

⁷Although we could have contrived a task with an objective ‘best’ answer—best model or prompt template—that wouldn’t reflect the kind of ambiguities present in many real world decisions around LLM usage.

of their investigations, not to guide them towards particular domains of interest. The one exception is where a participant only queried a single model; in this case, the researcher could suggest that the user try querying multiple models at once. Participants used the exact same interface as the public version of our tool, and had access to OpenAI’s *gpt-3.5* and *gpt-4*, Anthropic’s *claude-2*, Google’s *chat-bison-001*, and HuggingFace models.

After the tasks, we held a brief post-interview (5-10 min), asking participants to rate the interface (1-5) and explain their reasoning, what difficulties they encountered, suggestions for improvements, whether they felt their understanding of AI was affected or not, and whether they would use the interface again and why.

5.1 Recruitment, Participant Demographics, and Data Analysis

We recruited in-lab participants around our U.S.-based university through listservs, Slack channels, and flyers. We tried to expand our reach beyond people experienced in CS and ML, specifically targeting participants in humanities and education. Participants were generally in their twenties to early thirties (nine 23-27; eight 28-34; three 18-22; one 55-64), predominantly self-reported as male (14 men, 7 women), and largely had backgrounds in computing, engineering, or natural sciences (ten from CS, data science, or tech; seven from bioengineering, physics, material science, or robotics; two from education; one from medicine and one from design). They had a moderate amount of past experience with AI text generation models (mean=3.3, stdev=1.0); one had none. Past Python programming experience varied (mean=3.1, stdev=1.3), with less experience in JavaScript (mean=2.0, stdev=1.3); two had no programming experience. Eight had “worked on an academic study, paper, or project that involved evaluating large language models.” All participants came in to the lab, with studies divided equally among the first three coauthors. Each study took 75 minutes, and participants were given \$30 in compensation (USD). Due to ethical concerns surrounding the overuse of Amazon gift cards in human subject studies [27, 31], we paid all participants in cash.

For our interview study, we sought participants who had already used ChainForge for real-world tasks, reaching out via social media, GitHub, and academic networks. The first author held six semi-structured, 60 min. interviews with eight participants (in two interviews, two people had worked together). Interviews took place via videoconferencing. Interviewees were asked to share their screen and walk through something they had created with ChainForge. Unlike our in-lab study, we kept interviewees’ screen recordings private unless they allowed us to take a screenshot, since real-world users are often working with sensitive information. Interviewees generously volunteered their time.

We transcribed all 32 hours of screen recordings and interviews, adding notes to clarify participant actions and references (e.g., “[Opens inspector; scrolls to top]. It seems like it went fast enough... [Reading from first email group] ‘Hi...’”). We noted conceptual or usability problems and the content of participant references. We analyzed the transcripts through a combination of inductive thematic analysis through affinity diagramming, augmented with a spreadsheet to list participants’ ideas, behaviors (nodes added, process of their exploration, whether they imported data, etc), and answers

to post-interview questions. For our in-lab study, three coauthors separately affinity diagrammed three transcripts each, then met and joined the clusters through mutual discussion. The merged cluster was iteratively expanded with more participant data until clusters reached saturation. For interviews, the first author affinity diagrammed all transcripts to determine themes. In what follows, in-lab participants are P1, P2, etc.; interviewees are Q1, Q2, etc.

6 MODES OF PROMPT ENGINEERING AND LLM HYPOTHESIS TESTING

What process do people follow when prompt engineering and testing hypotheses about LLM behavior more generally? Before we break down findings per study, we provide a birds-eye view of how participants in general used ChainForge. Synthesizing across studies, we find that people tend to move from an *opportunistic exploration* mode, to a *limited evaluation* mode, to an *iterative refinement* mode. About half of our in-lab users, especially end-users with limited prior experience, never left exploration mode; while programmers or auditors of LLMs quickly moved into limited evaluation mode. Some interviewees had disconnected parts of their flows that corresponded to exploration mode, then would scroll down to reveal extensive evaluation pipeline(s), explaining they had transferred prompts from the exploratory part into their evaluation. In Appendix A, we provide one Case Study for each mode. Notice how these modes correspond to users moving from the left side of Fig. 2 towards the right.

Opportunistic exploration mode is characterized by rapid iteration on prompts, input data, and hypotheses; a limited number of prompts and input data; and multi-model comparison. Users *prompt / inspect / revise*: send off a few prompts, inspect results, revise prompts, inputs, hypotheses, and ideas. In this mode, users are sending off quick experiments to probe and poke at model behavior (“*throw things on the wall to see what’s gonna stick*”, Q3). For instance, participants who conducted adversarial attacks like jailbreaking [6] would opportunistically try different styles of jailbreak prompts, and were especially interested in checking which model(s) they could bypass.

Limited evaluation mode is characterized by moving from ad-hoc prompting to *prototyping an evaluation*. Users have reached the limits of manual inspection and now want a more efficient, “at-a-glance” test of LLM behavior, achieved by encoding criteria into automated evaluator(s) to score responses. Users *prompt / evaluate / visualize / revise*: prompt model(s), score responses downstream in their chain, visualize results, and revise their prompts, input data, models, and/or hypotheses accordingly. Hallmarks of this mode are users setting up an analysis pipeline, iterating on their evaluation itself, and “scaling up” input data. The evaluation is “limited” as evaluations at this stage are often “coarse”—for example, rather than checking factuality, check if the output is formatted correctly.⁸

⁸Crucially, this mode does *not* imply multiple prompts: a few in-lab participants set up an evaluation pipeline that *only sent off a single prompt*. Though these participants did add a TextFields or Tabular Data node, it only had a single value/field. Indications were that, with more time, they would have “scaled up” how many inputs or parameters they were sending to test more specific hypotheses. However, some users might have also had conceptual trouble imagining how to scale up their testing; we discuss this more later.

Iterative refinement mode is characterized by having an already-established evaluation pipeline and criteria and *tweaking* prompt templates and input data through further parametrization or direct edits, setting up one-off evaluations to check effects of tweaks, increasing input data complexity, and removing or swapping out models. Users *tweak / test / refine*: modify or parametrize some aspect of their pipeline, test how tweaks affect outputs compared to their “control”, and refine the pipeline accordingly. The key difference between limited evaluation and iterative refinement is in the solidity of the chain: here, users’ prompts, input data, and evaluation criteria have largely stabilized, and they are looking to *optimize* (e.g., through tweaks to their prompt, or extending input data to identify failure modes). Some interview participants had reached this mode, and were refining prompt templates or scaling up input data. The few in-lab participants that had brought in “prompt engineering” problems by importing prompt templates or spreadsheets would immediately set up evaluation pipelines, moving towards this mode.

These modes are suggestive and not rigidly linear; e.g., users may scrap their limited evaluation and return to opportunistic exploration. In Sections 7 and 8 below, we delve into specific findings for each study. For our in-lab study, we describe how people selected prompts and models, how ChainForge supports exploration and understanding, and note conceptual and usability issues. For our interview study, we focus on what differed from in-lab users.

7 IN-LAB STUDY FINDINGS

On average, participants rated the interface a 4.19/5.0 (stdev=0.66). No participant rated it lower than a three. When asked for a reason for their score, participants generally cited minor usability issues (e.g., finicky when connecting nodes, color palette, font choice, more plotting options). Eighteen participants wanted to use the interface again; five before being explicitly asked. Some just wanted to play around, citing model comparison and multi-response generation. Participants who had prior experience testing LLM behavior in academia or industry cited speed and efficiency of iteration as the primary value of the tool (“If I had started with using this, I’d have gotten much further with my prompt engineering... This is much faster than a Jupyter Notebook”, P4; “this would save me half a day for sure... You could do a lot of stuff with it”, P21). Participants mentioned prior behavior as having multiple tabs open to chat with different models, manually copying responses into spreadsheets, or writing programs. Three wanted to use ChainForge for research.

We recount participants’ behavior in the structured task to choose a model and prompt template, overview how ChainForge supported participants’ explorations and understanding, and reflect on pain points.

7.1 How People Decide on Models and Prompts

How do people choose a text generation model or prompt, when presented with side-by-side responses? People appear to weigh *trade-offs* in response quality for different criteria and contexts of usage. Participants would perceive one prompt or model to excel in one criteria or context, but do poorly in another; for another prompt or model, it was vice-versa. Here, we use “criteria” liberally

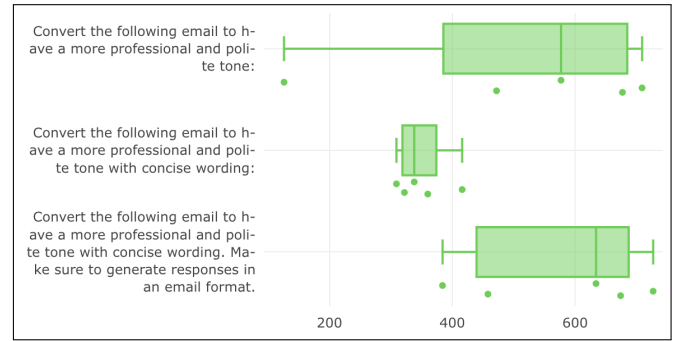


Figure 6: P17 plots the response lengths of three command prompts, augmenting her theories about each prompts’ performance.

to mean both our explicit criteria and also participants’ tacit preferences. Participants would also implicitly *rank* criteria, assigning more weight to some over others, and refer to friction between criterias (e.g., P2 “*prefer[red] professional over concise, because it [email] can be concise, but misconstrued*”). Moreover, seeing *multiple representations* of prompt performance, each of which better surfaced aspects of responses that corresponded to different criteria, could affect participants’ theorizing and decision-making. We unpack these findings here.

For the first part of our structured task, participants reached no consensus on which model performed “better”: eight chose PaLM2, seven GPT-4, and six Claude-2. There was no pattern in reasoning. Participants *did* notice similar features of each models’ response style, but *how* they valued that style differed. Some participants liked some models for the same reason others disliked them; for instance, P1 praised PaLM2 for its lengthy emails; while P17 chose GPT-4 because “*PaLM2 is too lengthy*.” Although we had deliberately designed our first criteria against the outputs of Claude (for its explanatory information around the email), some participants still preferred Claude, perceived its explanations as useful to their imagined users, or preferring its writing style. In the unstructured task, participants developing apps also mentioned exogenous factors such as pricing, access, and response time when comparing models.

How did people choose one prompt among multiple? Like when choosing models, participants appeared to weigh trade-offs between different criteria and contexts. Having multiple representations (e.g., plots of prompt performance) could especially give users a different “view” that augmented understanding and theorizing. P1 describes tensions between his and his users’ needs, referencing both manual inspection and a plot of response lengths by prompt:

“If I am a developer, I like this one [third prompt] because it will help me better to pass the output... But if they [users] have a chance to see this graph [Vis node], they would probably choose this one [second prompt] because it fits their needs and it’s more concise [box-and-whiskers plot has smallest median and lowest variability]... So I think it depends on the view.”

Multiple representations could also augment users' theorizing about prompting strategy. For instance, P17 had three command prompts, each iteration just tacking more formatting instructions onto the end of the default prompt (Figure 6). Comparing between her plot and Table Layout, she theorizes: *"After adding 'generate response in an email format' it made it lengthier... But if I don't say 'with concise wording'... sometimes it generates responses that are three paragraphs, for a really simple request. So I would [go with] the second instruction... [and its] the length difference [variance] is less."* Seeing that one prompt resulted in shorter or less variable responses could cause a participant to revise an earlier opinion. After noticing via the plot that his first command *"seem[s] more consistent"*, P4 wanted to mix features from it into his chosen prompt to improve the latter's concision, as he still preferred the latter's textual quality.

These observations suggest that systematic evaluations can contest fixation [45] caused by manual inspection. However, it also reveals users may need *multiple* representations, or they will make decisions biased by features that are easiest to spot in only one. With multiple, they can make decisions more confidently, mixing and matching parts of each prompt to progress towards an imagined ideal. The benefit of prompt comparison also underscores the importance of starting from a *variety* of prompts—similar to past work [45], many of our participants struggled to come up with a variety of prompts, with thirteen just perturbing our initial command prompt. We reflect on this more in our Discussion.

7.2 ChainForge Supported a Variety of Use Cases and Users

Participants brought in a variety of ideas to the unstructured task, ranging from auditing of LLM behavior to refining an established prompt used in production. We recount three participants' experiences as Case Studies in Appendix A, each corresponding to a Mode of Usage from Section 6. Seven participants evaluated model behavior given concrete criteria, with six importing prior data; ideas ranged from testing model's ability to understand program patch files, to classifying user attitudes in messaging logs. Nine audited models in opportunistic exploration mode, looking for biases or testing limits (e.g., asking undecidable questions like *"Does God exist?"*, P20). Of these users, four conducted adversarial attacks [6], seemingly influenced by popular culture about jailbreaking. P9 and P15, both with no programming experience, used the tool to audit behavior, the former comparing models' ability to generate culturally-appropriate stories about Native Alaskans. Others were interested in generating text for creative writing tasks like travel itineraries. Participants often searched the internet, such as cross-checking factual data, copying prompts from Reddit, or evaluating code in an online interpreter. Overall, nine participants imported data (six with spreadsheets) to use in their flow.

7.3 ChainForge Affected Participants' Understanding of AI Behavior or Practice

In the post-interview, fifteen participants said their understanding of AI was affected by their experience. Six were surprised by the performance of Claude-2 or PaLM2, feeling that, when confronted with direct comparisons to OpenAI models, they matched or exceeded the latter's performance. Five said that their *strategy* of

prompting or prompt engineering had changed (*"[Before], I wasn't doing these things efficiently... I [would] make minor modifications and rerun, and that would take hours... Here, since everything is laid out for me, I don't want to give up"*, P4). Others less experienced with AI models learned about general behavior. P16, who had never prompted an AI model before, *"realized that different models have completely different ways of understanding my prompts and hence responding, they also have a completely different style of response."* P15, covered in Case Study A.1, said she had lost *"trust"* in AI.

7.4 Challenges and Pain-points

Though many participants derived value from ChainForge, that is not to say their experience was frictionless. The majority of usability issues revolved around the flow UI, such as needing to move nodes around to make space, connecting nodes and deleting edges; others related to inconsistencies in the ordering of plotted variables, and wanting more control over colors and visualizations. Some participants also encountered conceptual issues, which sometimes indicate users getting used to the interface. The most common conceptual issue was learning how prompt templating worked, and especially, forgetting to declare input variables in Prompt Nodes. Once users learned how to template, however, the issue often disappeared (*"prompt variables... there's a bit of a learning curve, but I think it makes sense, the design choice"*, P13). Learning template variables seemed related to past programming expertise and not AI, suggesting users without any prior programming experience will need extra resources.⁹

Import to reflect on is that, in the lab, researchers were on-hand to guide users. Although users were the ones suggesting ideas—often highly domain-specific ones—researchers could help users with ways to implement them and overcome conceptual hurdles. Some end-users *and even a few users with substantial prior experience with AI models or programming with LLM APIs* appeared to have trouble *"scaling up,"* or systematizing, their evaluations. For example, P10 rated themselves as an expert in Python (5) and had conducted prior research on LLM image models. They set up an impressive evaluation, complete with a prompt template, Prompt Node, Chat Turn, Simple Evaluator and Vis nodes, but ultimately only sent off a single prompt to multiple models. We remark more on this behavior in Discussion.

8 INTERVIEWS WITH REAL-WORLD USERS

Our interview findings complement, but in important places diverge, from our in-lab studies. Like many in-lab participants, real-world users praised ChainForge's features and used it for goals we had designed for—like selecting models or prompt testing—however, some things real users cared about were hardly, if ever mentioned by in-lab participants. As we analyzed the data and compared it with our in-lab study, we realized that many user needs and pain-points revolve around the fact that they were using ChainForge to *prototype data processing pipelines*, a wider context that re-frames the tasks we had designed ChainForge to support as subtasks of a larger goal. Interviewees remarked most about *easing the export*

⁹For instance, P16, who had never prompted an AI model before, attributed her adeptness with templating to prior programming experience. By contrast, P9 had no programming experience and struggled; after the intro video he was *"a little overwhelmed"*, asking: *"what language am I in?"*

and sharing of data from ChainForge, adding *processor nodes*, the importance of the *Inspect Node* for sharing and rapid iteration, and the *open-source nature* of the project for their ability to adapt the code to their use case. We discuss these insights more below; but first, we provide an overall picture, reviewing similarities, use cases, and concrete value that real-world users derived from ChainForge.

We list interview participants in Table 2, with use cases and nodes used. The Outcome column suggests the actionable value that ChainForge provided. Note that Q1 and Q2’s primary use case was building on the source code to enable their HCI research project. All six users of the *interface* found it especially useful for prototyping and iterating on prompts and pipelines (e.g., Q5: “I see the use case for ChainForge as a very good prompt prototyping environment”). Usage reflected modes of *limited evaluation* and *iterative refinement*, with multiple participants describing a prompt/evaluate/visualize/revise loop: query LLM(s), evaluate responses and view the plot, then refine prompts or change models, until one reaches the desired results. For instance, Q3 described tweaking a prompt template until the LLM output in a consistent format, facilitated by maximizing 100% bars in a Vis Node across all input data. Some participants saw ChainForge as a rapid prototyping tool missing from the wider LLMops ecosystem, a tool they used “until I get to the point where I can actually write it into hard code” (Q4). Three appreciated how few nodes there were in ChainForge given its relative power, compared to other node-based interfaces (e.g., Q8: “It’s impressive. What you’re able to accomplish with so few”). They worried that adding too many new nodes would make the interface more daunting for new users. Q4 and Q7 found it more effective than Jupyter notebooks (Q7: “I enjoyed ChainForge... because I could run the whole workflow over and over again, and... in Jupyter, that was not easy”). In the rest of this section, we expand upon differences from our in-lab study.

8.1 Prototyping data processing pipelines

Five interviewees were using ChainForge not (only) for prompt engineering or model selection, but for *on-demand prototyping of data processing pipelines involving LLMs*. All imported data from spreadsheets, then would send off many parametrized prompts, iterate on their prompt templates and pipelines, and ultimately export data to share with others. Such users also used ChainForge for at least one of its intended design goals, but always in service of their larger data processing goal. For Q7, “the idea was to write a pipeline that... helps you with this whole process of data cleaning.” For him, ChainForge was ideal for “whenever you have a variety of prompts you want to use on something particular, like a data set. And you want to explore or investigate something.” Another user, Q3, would open his refined flow, edit one value, re-run it and then export the responses to a spreadsheet. Like other participants, he remarked on ChainForge’s combinatorial power as its chief benefit, compared to other tools (“This tool is strong at prompt refining. With [Flowise]...Let’s say I wanted to try multiple [input fields]. I don’t think I could do that”). Participants also mentioned iterating on the *input data* as part of the prototyping process. Finally, related to data processing, three users wished for processor nodes, like “join” nodes to concatenate LLM responses, and in one case were manually copying LLM outputs into a separate flow to emulate

concatenation. Note that many needs and pain-points below are related to data processing.

8.2 Getting data out and sharing with others

Many participants wanted to export data out of ChainForge. This was also the most common pain point, especially when transitioning from a prototyping stage—which they perceived as ChainForge’s strong point—to a production stage (e.g., “it would be helpful when we are out of this prototyping stage, that the burden or the gap—changing the environment... gets tightened”, Q5). Needs broke down into two categories: exporting for integration into another application, and exporting for sharing results with others. For the former, developer users would use ChainForge to battle-test prompts, model behavior, and/or prompt chains, but then wished for an easier way to export their flows to text files or app building environments.¹⁰ For the latter, five interviewees shared results with others, whether through files, screenshots of their flows, exported Excel spreadsheets of responses, or copied responses. Q5 and Q6 stressed the importance of the *Inspect Node*—a node that no in-lab participant used or mentioned (“[Once] the result is worth documenting, you create an Inspect node.”). They took screenshots of flows and sent them to clients, in one case convincing a client to move forward with a project. The anticipation of sharing with others also could change behavior. Q3 had several TextFields nodes with only a single value, “because I knew that it was something that essentially other teams might want to change.” Sharing could also be a pain-point, with two wanting easier shareable “reports” of their analysis results.

8.3 Pain points: Hidden affordances and friction during opportunistic exploration mode

Like in-lab participants, interviewees also encountered usability and conceptual issues. A common theme was individual users expressing a need for features that already exist but are relatively hidden, surfaced only through examples or documentation. These hidden affordances included *implicit template variables*, *metavariables*, and *template chaining*. The former two features address users’ need to reference upstream metadata—metadata associated with input data or responses—further downstream in a chain.¹¹ Another pain point was friction during the opportunistic exploration phase. In Section 6, we mentioned some interviewees had disconnected regions of their flows, with one region we termed opportunistic exploration mode (rapid, early-stage iteration through input data, prompts, models, and hypotheses; usually, a chain of three nodes, TextField-Prompt-Inspect). In this mode, some interviewees preferred to inspect responses directly on the flow with an Inspect Node (instead of the pop-up window), as it facilitated rapid iteration.

¹⁰This does not, however, mean they perceived ChainForge as an “app building” tool—some even expressed worry about it trying to do too much. When asked about how he would feel if ChainForge supported app development, Q3 remarked, “I just worry about it [ChainForge] becoming too complicated. Like, are you building the app side of it?” He said even if ChainForge supported app-building, it would need different “modes”: “app building mode or prompt refining mode.”

¹¹For example, in ChainForge one can define a column of a table and then refer to it downstream via a metavariable (Fig. 5). However, Q5 and Q6 seemed unaware of this feature and had implemented a workaround. For implicit template variables, Q4 needed to reference an upstream value {gender} later downstream in a prompt chain, and was unaware that an implicit template variable could accomplish this (e.g. {#gender}).

ID(s)	Location	Work	How discovered	Use Case(s)	Major nodes and features used	Outcome
Q1-2	U.S. (west)	Acad.	Medium post	Adapting code for LLM image model prototyping	<i>Adapted source code (esp. nodes, model querying, cacheing)</i>	Submitting HCI paper to major conference in time for deadline
Q3	U.K.	Ind.	GitHub	Supporting requirements analysis in software testing	Prompt, Tabular Data, TextFields, Vis, JS Eval; export to excel, metavariables	<i>"I've had good conversations with [other] developers as a result"</i>
Q4	U.S. (east)	Acad.	Hacker-News	Auditing models for gender bias	Prompt, TextFields, Chat Turn, Inspect node; prompt chaining, LLM scoring	Improved prompt templates for a pipeline they are building in Python
Q5-6	Germany	Ind.	LearnPrompting.org	Information synthesis and extraction from documents	Prompt, Tabular Data, JS Eval, Vis, Comment, Inspect node; comparing models	<i>"Convinced client [to] continue with the next phase of [a] project"</i>
Q7	U.S., Austria	Acad.	Word of mouth	Cleaning a column of tabular data	Prompt, Tabular Data, TextFields, JS Eval, Vis; variables in code eval	Decided that LLMs were not reliable enough for their task
Q8	U.S. (central)	Acad.	Twitter	Extracting and formatting info from podcast episode metadata	Prompt, TextFields, CSV, JS Eval, Vis, Comment; comparing models, prompt chaining, template chaining	Found that neither OpenAI model produced good enough outputs; now looking into local models

Table 2: Our interviewees. All interviewees imported data into ChainForge, with the exception of Q1-2 (who adapted source code). Interviewees commonly had multiple evaluation branches in their flows, as well as multiple flows in the same document.

They wanted an even more immediate, in-context way to read LLM responses that would not require them to attach another node.¹²

8.4 Open-source flexibility

Multiple interviewees mentioned looking at our source code, and two projects extended it. Q5 and Q6, employees of a consulting firm that works with the German government, extended the code to support a German-based LLM provider, AlephAlpha [1], complete with a settings screen. They cited the value of supporting European businesses and GDPR data protection laws: *"the government [of Germany] wants to support it. It's a local player... [and] There's a strong need to to hide and to to protect your data. I mean, GDPR, it's very strict in this."* Their goal was to use ChainForge to determine "if it makes sense to switch to" the German model for their use cases, over OpenAI models. HCI researchers Q1 and Q2's chief interaction with the tool was its source code, finding it helpful for jumpstarting a project on a flow-based tool for LLM image model prototyping. Q2 appreciated the *"thought put into"* caching, Prompt Node progress bar, and multi-model querying, adding: *"It was very easy for me to set up ChainForge... [and it was] surprisingly easy to [extend]... a lot easier than I had expected."* They said that the jump-start ChainForge provided was a chief reason they were able to complete their project in time to submit a paper to the annual CHI conference.

9 DISCUSSION AND CONCLUSION

Our observations suggest that ChainForge is useful both in itself, but also as an 'enabling' contribution, an open-source project which others can extend (and are extending) to investigate their own ideas and topics, including other research publications to this very conference. Given that ChainForge was released only a few months ago, we believe the stories presented here provide evidence for its real-world usefulness. In what follows, we review our key findings.

¹²This need was again reflected in a later GitHub Issue and has since been addressed with a pull-out inspector drawer.

Our work represents one of the only "prompt engineering" system contributions with data about real-world usage, as opposed to in-lab studies on structured tasks. Some of what real users cared about, like features for exporting data and sharing, were absent from our in-lab study—and are, in fact, also absent from similar LLM-prompting-system research with in-lab studies [5, 25, 43–45]. Most surprising (to us) was that some knowledge workers were using ChainForge for a task we had never anticipated—*data processing*. Although we only had six interface users in our interview study, the only two in-lab participants in startups, P8 and P4, were both testing LLMs' ability to process and reformat data. Most prior LLM tools target sensemaking [15, 36], prompt engineering [14, 25], or app building [43], but do not specifically target, or even mention, data processing. Our findings suggest a need for systems to support *on-demand creation of data processing pipelines involving LLMs*, where the purpose is not (always) to make apps, but simply process data and share the results. ChainForge's combinatorial power—the ability to send off many queries at once, parametrized by imported data—appeared key to supporting this need. Future systems should go further by providing users more accessible ways to reference upstream metadata further downstream in their chain (see 8.3).

Second, we identified three modes of prompt engineering and LLM hypothesis testing: *opportunistic exploration*, *limited evaluation*, and *iterative refinement*. The first mode is similar to Barke et al.'s exploration mode for GitHub CoPilot [2]. Future systems should explicitly consider these modes when designing and framing the work. For instance, users often too quickly enter iterative refinement mode—refining on the first prompt they try—rather than exploring a variety before settling on one [45]. If a prompt engineering tool only targets iterative refinement, then the opportunistic exploration stage—finding a good prompt to begin with—may be too quickly skirted over, trapping users in potentially suboptimal prompting strategies. These modes also suggest design opportunities. For instance, we believe that ChainForge's design could have better supported opportunistic exploration mode, with some users

wanting a simpler way to inspect LLM responses in-context (8.3). One design solution may be to concretize each mode into separate, related interfaces or layouts—e.g., a more chat-like interface for exploration mode, that then facilitates the transition to later modes, each with dedicated interfaces. Prior LLM-prompting systems seem to target opportunistic exploration [15, 36] or iterative refinement [25, 35], but overlook *limited evaluation*: an important mid-way point characterized by prototyping small-scale, quick-and-messy evaluations on the way to greater understanding. Future work might target the prototyping of on-demand LLM evaluation pipelines themselves (see “model sketching” for inspiration [18]).

Third, we found that when people choose different prompts and models, they weigh *trade-offs* in performance for different criteria and contexts, and bring their own perspectives, values, preferences, and contexts to bear on decision-making. Having multiple representations of responses seemed to help participants weigh trade-offs, rank prompts and models, develop better mental models, and make revisions to their prompts or hypotheses more confidently. Connecting to theories of human learning [10, 23], the case study in A.1 suggests that cross-model comparison might also help novices improve mental models of AI by forcing them to encounter differences in factual information, jarring AI over-reliance [20]. The subjectivity of choosing a model and prompt implies that, while LLMs can certainly *help* users generate or evaluate prompts [5, 46], there will never be such a thing as *fully* automated prompt engineering. Rather than framing prompt engineering (purely) as an optimization problem, projects looking to support prompt engineering should instead look for ways to give users greater *control* over their search process (e.g., “steering” [5, 46]).

A final point and caveat: while users found ChainForge useful for *implementation* and *iteration*, including on real-world tasks, more work needs to be done on *conceptualization* and *planning* aspects, to help users move out of opportunistic exploration into more systematic evaluations. In-lab users seemed limited in their ability to imagine systematizing their tests, *even a few with prior expertise in AI or programming with LLM APIs*. This extends prior work studying how “non-AI-experts” prompt LLMs [45], suggesting even people who otherwise perceive themselves to be AI experts may have trouble systematizing their evaluations. Since LLMs are nondeterministic (at least, often queried at non-zero temperatures) and prone to unexpected jumps in behavior from small perturbations, it is important that future systems and resources help reduce fixation and guide users from early exploration into systematic evaluations. We might leverage concepts from tools designed for more targeted use cases; e.g., the auditing tool AdaTest++ provides users “prompt templates that translate experts’ auditing strategies into reusable prompts” [33, p. 15–6]. Other work supports creation of prompts or searching of a “prompt space” [25, 34, 35]. To support systematization/scaling up, we might also employ an interaction whereby a user chats with an AI that sketches out an evaluation strategy.

9.1 Limitations

Our choice to use a qualitative evaluation methodology derived from well-known difficulties around toolkit research [19, 28], concerns about ecological validity, and, most importantly, from the

fact that we could not find a prior, well-established interface that matched the entire featureset of ChainForge. Our goal was thus to establish a baseline system that future work might improve upon. While we believe our qualitative evaluation yielded some important findings, more quantitative, controlled approaches should be performed on parts of the ChainForge interface to answer targeted scientific questions. Our in-lab study was also of a relatively short duration (75 min); future work might observe changes in user behavior over longer timeframes, for instance with a multi-week workshop. Finally, for our interview study, we acknowledge a self-selection bias, where participating interviewees may already have found ChainForge useful, missing users who did not. Our in-lab study provided some insights—we speculate that users’ prior exposure to programming was important to the quality of their experience.

ACKNOWLEDGMENTS

This work was partially funded by the NSF grants IIS-2107391, IIS-2040880, and IIS-1955699. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

REFERENCES

- [1] Aleph-Alpha. 2023. Aleph-Alpha. <https://www.aleph-alpha.com> Accessed: Sep 2 2023.
- [2] Shraddha Barke, Michael B James, and Nadia Polikarpova. 2023. Grounded copilot: How programmers interact with code-generating models. *Proceedings of the ACM on Programming Languages* 7, OOPSLA1 (2023), 85–111.
- [3] Luca Beurer-Kellner, Marc Fischer, and Martin Vechev. 2023. Prompting is programming: A query language for large language models. *Proceedings of the ACM on Programming Languages* 7, PLDI (2023), 1946–1969.
- [4] Markus Binder, Bernd Heinrich, Marcus Hopf, and Alexander Schiller. 2022. Global reconstruction of language models with linguistic rules—Explainable AI for online consumer reviews. *Electronic Markets* 32, 4 (2022), 2123–2138.
- [5] Stephen Brade, Bryan Wang, Mauricio Sousa, Sageev Oore, and Tovi Grossman. 2023. Promptify: Text-to-Image Generation through Interactive Prompt Exploration with Large Language Models. *arXiv preprint arXiv:2304.09337* (2023).
- [6] Gelei Deng, Yi Liu, Yuekang Li, Kailong Wang, Ying Zhang, Zefeng Li, Haoyu Wang, Tianwei Zhang, and Yang Liu. 2023. Jailbreaker: Automated Jailbreak Across Multiple Large Language Model Chatbots. *arXiv preprint arXiv:2307.08715* (2023).
- [7] Harrison Chase et al. 2023. LangChain. <https://pypi.org/project/langchain/>.
- [8] FlowiseAI, Inc. 2023. FlowiseAI Build LLMs Apps Easily. flowiseai.com.
- [9] Nat Friedman, Zain Huda, and Alex Lourenco. 2023. Nat.Dev. <https://nat.dev/>.
- [10] Dedre Gentner and Arthur B Markman. 1997. Structure mapping in analogy and similarity. *American psychologist* 52, 1 (1997), 45.
- [11] 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>
- [12] Chip Huyen. 2022. *Designing machine learning systems*. O'Reilly Media, Inc.
- [13] Jest. 2023. Jest. <https://jestjs.io/>.
- [14] Ellen Jiang, Kristen Olson, Edwin Toh, Alejandra Molina, Aaron Donsbach, Michael Terry, and Carrie J Cai. 2022. PromptMaker: Prompt-based Prototyping with Large Language Models. 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 35, 8 pages. <https://doi.org/10.1145/3491101.3503564>
- [15] 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>
- [16] Laewoo (Leo) Kang, Steven J. Jackson, and Phoebe Sengers. 2018. Intermodulation: Improvisation and Collaborative Art Practice for HCI. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems* (Montreal, QC, Canada) (CHI '18). Association for Computing Machinery, New York, NY, USA, 1–13.

- <https://doi.org/10.1145/3173574.3173734>
- [17] 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>
 - [18] Michelle S. Lam, Zixian Ma, Anne Li, Izequiel Freitas, Dakuo Wang, James A. Landay, and Michael S. Bernstein. 2023. Model Sketching: Centering Concepts in Early-Stage Machine Learning Model Design. 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 741, 24 pages. <https://doi.org/10.1145/3544548.3581290>
 - [19] David Ledo, Steven Houben, Jo Vermeulen, Nicolai Marquardt, Lora Oehlberg, and Saul Greenberg. 2018. Evaluation Strategies for HCI Toolkit Research. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems* (Montreal, QC, Canada) (CHI '18). Association for Computing Machinery, New York, NY, USA, 1–17. <https://doi.org/10.1145/3173574.3173610>
 - [20] Q. Vera Liao and S. Shyam Sundar. 2022. Designing for Responsible Trust in AI Systems: A Communication Perspective. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency* (Seoul, Republic of Korea) (FAccT '22). Association for Computing Machinery, New York, NY, USA, 1257–1268. <https://doi.org/10.1145/3531146.3533182>
 - [21] Mark Liffiton, Brad Sheese, Jaromir Savelka, and Paul Denny. 2023. CodeHelp: Using Large Language Models with Guardrails for Scalable Support in Programming Classes. *arXiv preprint arXiv:2308.06921* (2023).
 - [22] Logspace. 2023. LangFlow. <https://www.langflow.org/>.
 - [23] Ference Marton. 2014. *Necessary conditions of learning*. Routledge.
 - [24] Microsoft. 2023. Prompt Flow. <https://microsoft.github.io/promptflow/>.
 - [25] Aditi Mishra, Utkarsh Soni, Anjana Arunkumar, Jinbin Huang, Bum Chul Kwon, and Chris Bryan. 2023. PromptAid: Prompt Exploration, Perturbation, Testing and Iteration using Visual Analytics for Large Language Models. *arXiv preprint arXiv:2304.01964* (2023).
 - [26] Graham Neubig and Zhiwei He. 2023. Zeno GPT Machine Translation Report.
 - [27] Wing Ng, Ava Anjom, and Joanna M Drinane. 2022. Beyond Amazon: Social Justice and Ethical Considerations for Research Compensation. *Psychotherapy Bulletin* (2022), 17.
 - [28] Dan R. Olsen. 2007. Evaluating user interface systems research. In *Proceedings of the 20th Annual ACM Symposium on User Interface Software and Technology* (Newport, Rhode Island, USA) (UIST '07). Association for Computing Machinery, New York, NY, USA, 251–258. <https://doi.org/10.1145/1294211.1294256>
 - [29] OpenAI. 2023. OpenAI Playground. <https://platform.openai.com/playground>.
 - [30] OpenAI. 2023. openai/evals. <https://github.com/openai/evals>.
 - [31] Jessica Pater, Amanda Coupe, Rachel Pfafman, Chanda Phelan, Tammy Toscos, and Maia Jacobs. 2021. Standardizing Reporting of Participant Compensation in HCI: A Systematic Literature Review and Recommendations for the Field. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (Yokohama, Japan) (CHI '21). Association for Computing Machinery, New York, NY, USA, Article 141, 16 pages. <https://doi.org/10.1145/3411764.3445734>
 - [32] Fábio Perez and Ian Ribeiro. 2022. Ignore Previous Prompt: Attack Techniques For Language Models. In *NeurIPS ML Safety Workshop*.
 - [33] Charvi Rastogi, Marco Tulio Ribeiro, Nicholas King, and Saleema Amershi. 2023. Supporting Human-AI Collaboration in Auditing LLMs with LLMs. *arXiv preprint arXiv:2304.09991* (2023).
 - [34] Fobo Shi, Peijun Qing, Dong Yang, Nan Wang, Youbo Lei, Haonan Lu, and Xiaodong Lin. 2023. Prompt Space Optimizing Few-shot Reasoning Success with Large Language Models. *arXiv preprint arXiv:2306.03799* (2023).
 - [35] Hendrik Strobelt, Albert Webson, Victor Sanh, Benjamin Hoover, Johanna Beyer, Hanspeter Pfister, and Alexander M Rush. 2022. Interactive and visual prompt engineering for ad-hoc task adaptation with large language models. *IEEE transactions on visualization and computer graphics* 29, 1 (2022), 1146–1156.
 - [36] Sangho Suh, Bryan Min, Srishti Palani, and Haijun Xia. 2023. Sensecape: Enabling Multilevel Exploration and Sensemaking with Large Language Models. *arXiv preprint arXiv:2305.11483* (2023).
 - [37] Tianxiang Sun, Yunfan Shao, Hong Qian, Xuanjing Huang, and Xipeng Qiu. 2022. Black-box tuning for language-model-as-a-service. In *International Conference on Machine Learning*. PMLR, 20841–20855.
 - [38] TruLens. 2023. trulens: Evaluate and Track LLM Applications. <https://www.trulens.org/>.
 - [39] Vellum. 2023. Vellum The dev platform for production LLM apps. <https://www.vellum.ai/>.
 - [40] Vercel. 2023. Vercel: Deveop.Preview.Ship. <https://vercel.com/>.
 - [41] Ian Webster. 2023. promptfoo: Test your prompts. <https://www.promptfoo.dev/>.
 - [42] Weights and Biases. 2023. Weights and Biases Docs: Prompts for LLMs. <https://docs.wandb.ai/guides/prompts>.
 - [43] 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>
 - [44] 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>
 - [45] 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>
 - [46] Yongchao Zhou, Andrei Ioan Muresanu, Ziwen Han, Keiran Paster, Silviu Pitit, Harris Chan, and Jimmy Ba. 2022. Large language models are human-level prompt engineers. *arXiv preprint arXiv:2211.01910* (2022).

A CASE STUDIES FOR MODES OF USAGE

To help readers understand how people used ChainForge and how their interactions varied, we walk through three participants' experiences. Each Case Study illustrates one mode from Section 6.

A.1 Opportunistic exploration mode: Iterating on hypotheses through rapid discovery of model behavior.

A graduate student from Indonesia, P15 wanted to test how well AI models knew the Indonesian education participation rate and could give advice on “*what the future of us as educators need to do.*” She opens a browser tab with official data from Badan Pusat Statistik (BDS), Indonesia’s Central Agency for Statistics. She wants to know “*what is the difference, if I use a different language?*” She adds a TextFields with two fields, one prompt in English, “*Tell me the participation rate of Indonesian students going to university*”; the second its Indonesian translation. “*Let’s just try two. I just want to see where it goes.*” Collecting responses, she looks over side-by-side responses of three models to her English prompt. All models provide different years and percentages. Scrolling down and expanding the response group for her Indonesian prompt, she finds that Falcon.7B only repeats her prompt and the PaLM2 model has triggered a safety filter.¹³ The last model, GPT-3.5, gives a different statistic than its English response.

Looking over these responses in less than a minute, P9 has discovered three aspects of AI model behavior: first, that models differ in their “facts”; second, that some models can refuse to answer when queried in a non-English language; third, that the *same* models can differ in facts when queried in a different language. She compares each number to the BDS statistics, finding them inaccurate. “*Oh my god, I’m curious. Why do they have like different answers across [models]?*” She then adds models to the Prompt Node. “*Can I try all [models]? I want to see if it’s in the table.*”

She queries the new models. A new hypothesis brews: “*In our prompt, [do] we need to say our source of data? Would that be like, more accurate?*” She wonders if different models are pulling data from different sources. Inspecting responses, she finds some models have cited sources of data: Claude cites UNESCO and GPT-4 cites the World Bank, UNESCO, and the Indonesian Ministry of Education

¹³This is a real problem with PaLM2 that we have communicated with the Google AI team about; they identified the issue and are fixing it.

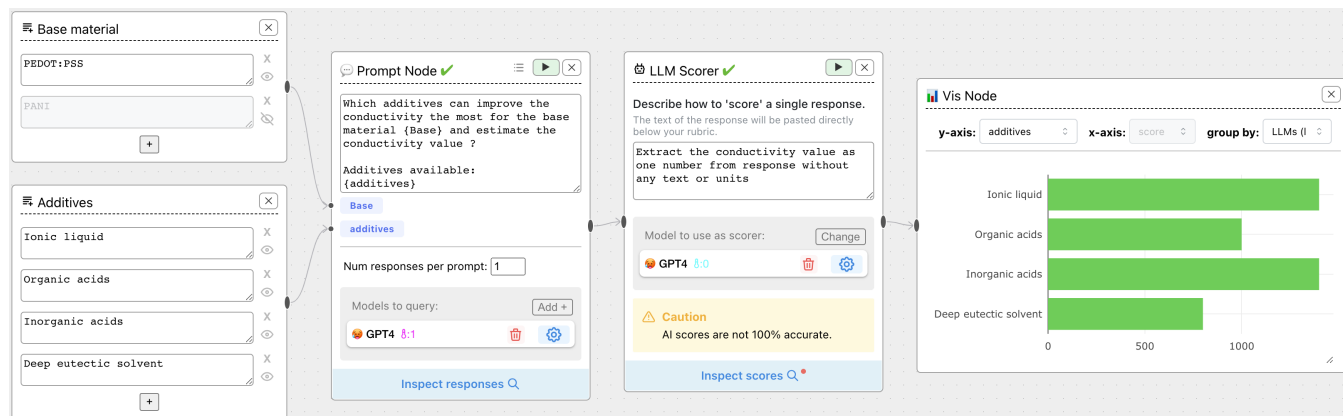


Figure 7: P18’s final flow, with one value toggled off to display their initial plot. The user asked GPT-4 to estimate how much the conductivity to a base polymer, PEDOT:PSS, may increase given one of four additives. They use an LLM Scorer to extract the mentioned conductivity value; after spot-checking it using Inspect, they plot in a Vis Node, finding it “roughly” correct.

and Culture. For her Indonesian prompt, she discovers that the same models only cite BPS in their responses. “BPS is only mentioned when I use Indonesian... For the English [prompt]... [it’s] more like, global... Wow, it’s very interesting how, the different language you use, there’s also a different source of data.”

She adds a second prompt variable, {organization}, to her prompt template. She attaches values World Bank, UNESCO, and Badan Pusat Statistik to it.¹⁴ Re-sending queries and inspecting responses, she expands the subgroups for BPS under both her Indonesian and English response groups, such that the two subgroups are on the same screen. When asking for BPS data in English, both GPT-3.5 and Claude refuse to answer, whereas the same models provide BPS numbers when asked in Indonesian. Moreover, Claude’s English response suggests the reader look at World Bank and UNESCO data instead, citing those sources. “That’s really interesting. Wow.”

Although the study ended here, this case illustrates hypothesis iteration, limited prompts, and eagerness for cross-model comparisons, key aspects of opportunistic exploration mode. With more time, the user might have set up an evaluation to check how models cite “global” sources of information when queried in English, compared to Indonesian.

A.2 Limited evaluation mode: Setting up an evaluation pipeline to spot-check factual accuracy.

How do users transition from exploratory to limited evaluation mode? We illustrate prototyping an evaluation and “scaling up” with P18, a material design student who used ChainForge to check an LLM’s understanding conductivity values of additives to polymers. The example also depicts a usability issue as the user scaled up.

Like Case #1, P18 begins in Opportunistic Exploration mode. They *prompt / inspect / refine*—send off queries, inspect responses, revise input data or prompts. They create a prompt template with two variables: *Base* and *additives* (Fig. 7). Initially they start with only one *Base*, and four *additives*. Inspecting responses, P18 is

impressed with GPT-4’s ability to suggest and explain *specific* additives under P18’s broad categories (e.g., EMIMBF4 for *Ionic Liquid*). They refine their questioning: “I want to estimate the approximate conductivity value.” They amend their prompt template, adding “and estimate the conductivity value”. Reviewing responses, they find the numeric ranges roughly correct.

They then wish to inspect the numbers in a more systematic fashion than manual inspection, and move into Limited Evaluation mode. The researcher helps P18 with how to extract the numbers, using an evaluator node, LLM Scorer, which they only saw once in the intro video. With this node, users can enter a natural language prompt to score responses. P18 iterates on the scorer prompt through a prompt/inspect/refine loop: first asking just for the number, then adding “without units” after they find it sometimes outputs units.¹⁵ “This is good. So we add some Vis Node.” They plot by *additive* on the y-axis (Fig. 7). “Very good. [Researcher: Is this true?] Roughly, yes. Roughly.”

P18 then wants to “scale up” by adding a second polymer to their *Base* variable. They search Google for the abbreviation of a conducting polymer, Polyaniline (PANI). They paste it as a second field and re-query the prompt and scorer nodes. Skimming scores in Table Layout in two seconds: “Oh, wow... It’s really good. Because PEDOT is most [conductive].” Inspecting the Vis Node, they encounter a usability limitation: they want to *group by Base*, when *additive* is plotted in y-axis, but cannot. Plotting by *Base* on the y-axis, they see via box-and-whiskers plot that PANI is collectively lower than PEDOT. They ask the researcher to export the evaluation scores.

This example illustrates limited evaluation mode, such as iterating on an evaluation pipeline (refining a scoring prompt), and beginning to “scale up” by extending the input data after the pipeline is set up. The user also encountered friction with usability when scaling up, wanting more options for visualization as input data complexity increased.

¹⁴Note that this part is in English now for both queries.

¹⁵Here is where a framework like LMQL [3] may come in handy to improve usability of this node.

A.3 Iterative Refinement mode: Tweaking an established prompt and model to attempt an optimization.

P8 works with a German startup, and brought in a prompt engineering problem, importing a dataset and prompt template from LangChain [7] (“we’re building a custom LLM app for an e-commerce company, a virtual shop assistant”). This template had already underwent substantial revisions; thus, the participant immediately moved into iterative refinement mode, allowing us to observe interactions we could only glimpse retroactively in our interviews.

P8’s startup was using GPT-4 (because “GPT-3.5 in German is really not that good”), but was curious about whether other models could perform better. He knew of Claude and PaLM2, but had been put off by needing to code up custom API calls. He also had a hypothesis that using English in parts of his German prompt would yield better results. Upon entering the unstructured task, he imported a spreadsheet with a Tabular Data Node and pasted his three-variable prompt template in a Prompt Node, connecting them up. He then added a Python Evaluator Node to check whether the LLM stuck to a length constraint he had put in his template. Using Grouped List layout, he compared responses between Claude and GPT-4 across ten input values for variable *product_information*. “GPT4 is going over [too long]... Claude seems to be fairly good at sticking—[opens another response group], actually, you know, we have an outlier here.”

Looking over responses manually, he implies that he had been manually evaluating each response (prior to the study) across his ten criteria. “I gave it... almost 10 instructions... Formal language, length, and so on. And for each... I now need to review it.” He notices that one of Claude’s responses includes the word *Begleiter*, a word he had explicitly instructed it to exclude: “Because that was a pattern I noticed with GPT-4 that it kept using this word... So I’m going to try now... how is Claude behaving if I give this instruction in English, rather than [German]?”

To test this, he abstracts the “avoid the following words” part of his prompt template into a new variable, *{avoid_words_instruction}*. He pastes the previous command into a TextFields, and add a second one—the same command but in English. He adds a Simple Evaluator node, checking if the response contains *Begleiter*. In Grouped List layout, he groups responses by *avoid_words_instruction* and click “Only show scores” to only see true/false values (**false** in red). Glancing: “So it’s not very statistically significant. But... GPT-4 never made the mistake, and Claude made the mistake with both English and German... So it doesn’t matter which language... [Claude] will still violate the instructions.” He attaches another Simple Evaluator to test another term, remarking that in practice he would write a Python script to test all cases at once, but the study is running out of time. “So Claude again violates it in both cases... [But for] English, it only violates it once—again—and in German it violates it twice. So maybe it’s slowly becoming statistically significant.” As the study ends, he declares that his investigation justified his original choice: “I should probably keep using GPT-4.”

Here we see aspects of iterative refinement mode—the participant has already optimized their prompt (pipeline) and is trying to tweak the prompt and model to see if they can improve the outputs even further, according to specific criteria. As we found in our structured

task, in making decisions, users weigh trade-offs between how different models and/or prompts fulfill specific criteria, and also rank criteria importance. For P8, his “avoid-words” criteria seemed mission-critical, whereas word count—which he perceived Claude better at sticking to—was evidently less important.

B LIST OF NODES

Node Name	Usage	Special Features
Inputs		
TextFields Node	Specify input values to a template variables in prompt or chat nodes.	Supports templating; can declare variables in brackets {} to chain inputs together.
CSV Node	Specify input data as comma-separated values. Good for specifying many short values.	Brackets {} in data are escaped by default.
Tabular Data Node	Import or create a spreadsheet of data to use as input to prompt or chat nodes.	Output values “carry together” when filling in multiple variables in a prompt template. Brackets {} in data are escaped by default.
Generators		
Prompt Node	Prompt one or multiple LLMs. Declare template variables in {}s to attach input data.	Can chain together. Can set number of generations per prompt to greater than one.
Chat Turn Node	Continue a turn of conversation with one or multiple chat LLMs. Supports templating of follow-up message.	Attach past prompt or chat output as context. Can also change what LLM(s) to use to continue conversation.
Evaluators		
JavaScript Evaluator	Write a JavaScript function to ‘score’ a single response. Scores annotate responses.	Boolean ‘false’ values display in red in response inspector. <i>console.log()</i> prints to node.
Python Evaluator	Same as JavaScript Evaluator but for Python.	Can import packages and use <i>print()</i> .
LLM Scorer	Prompt an LLM to score responses. (GPT-4 at zero temperature is default.)	Unlike prompt chaining, this attaches the scores as annotations on existing responses.
Simple Evaluator	Specify simple criteria to score responses as true if they meet the criteria.	Can test whether response <i>contains</i> , <i>starts with</i> , etc. a certain value. Can also compare against prompt variables or metavariables.
Visualizers		
Vis Node	Plot evaluation scores. Currently only supports boolean and numeric scores.	Plots by LLM by default. Change y-axis to plot by different prompt variables.
Inspect Node	Inspect LLM responses like the pop-up inspector, only inside a flow.	Only supports Grouped List layout.

Table 3: The nodes in ChainForge, grouped by type. A final node, the “Comment Node”, allows users to write comments.