

Prompt Engineering as a New Literacy: The Linguistics of Talking to Large Language Models

Saad Rehman Babary

Senior Software Engineer, Computer Science, University of Engineering and Technology, Lahore, Pakistan

Corresponding author: Saad Rehman Babary (Email: saadbabary97@gmail.com)

Received: 30/02/2026, Revised: 27/04/2026, Accepted: 06/05/2026

ABSTRACT

Since the release of large-scale generative language models, the practice of writing prompts, that is, the natural-language instructions given to a model to elicit a desired output, has evolved from an engineering curiosity into a widely discussed skill, variously labelled prompt engineering, AI literacy, or prompt literacy. This paper argues that prompting is best understood not merely as a technical or engineering competence but as a linguistic practice, one that draws on pragmatics, speech-act theory, register variation, and politeness strategies in ways that closely parallel ordinary human-to-human communication while also departing from it in theoretically interesting ways. Drawing on scholarship published between 2019 and 2026, spanning the original few-shot learning paradigm, systematic surveys of prompting techniques, Gricean and speech-act analyses of human-AI dialogue, and the emerging educational literature on prompt and AI literacy, the study develops a linguistic-functional taxonomy of prompting strategies and applies it, through qualitative content analysis, to a purposively selected corpus of widely documented prompt patterns, including zero-shot and few-shot instructions, chain-of-thought reasoning prompts, role or persona prompts, and structured academic frameworks such as the CLEAR model. The analysis finds that effective prompts systematically deploy directive speech acts, explicit and implicit context-setting, register calibration, and cooperative-principle-like maxims of quantity, quality, relation, and manner, adapted to the specific inferential behaviour of large language models. The study concludes that prompt engineering constitutes a genuine new literacy in the sociolinguistic sense, a describable, teachable set of language practices for achieving communicative goals with a non-human interlocutor, and it proposes an integrated pedagogical and linguistic research agenda for studying this literacy as it continues to develop.

Keywords: Prompt Engineering, Large Language Models, Linguistics, Pragmatics, AI Literacy, Prompt Literacy, Human-AI Interaction, Speech Acts

1. Introduction

For most of the history of human-computer interaction, communicating with a machine meant learning the machine's language: a rigid syntax of commands, menus, or query strings that demanded precision and offered no tolerance for ambiguity, incompleteness, or implicature. The emergence of large language models trained to complete and respond to natural-language text has inverted this relationship. Users no longer need to learn a formal command language; instead, they address the system in ordinary English, or another natural language, and the system attempts to infer their intent from the wording, structure, and context of that address. This shift has given rise to a new communicative skill, widely termed prompt engineering, the deliberate construction of natural-language input, or prompts, designed to elicit a specific, high-quality response from a large language model (Liu, Yuan, Fu, Jiang, Hayashi, & Neubig, 2023).



The technical literature on prompt engineering has grown rapidly since the demonstration, in the GPT-3 paper by Brown and colleagues (2020), that large language models could perform new tasks from a handful of natural-language examples embedded directly in the input, without any additional training. This finding reframed the central bottleneck of working with such models: performance depended less on the model's underlying parameters, which the ordinary user cannot change, and more on the wording of the instruction given to it, a variable entirely under the user's control. Since then, an extensive body of research has catalogued prompting techniques, from simple zero-shot instructions to elaborate chain-of-thought and role-based prompts, and has increasingly framed the ability to construct such prompts not as a narrow technical skill but as a form of literacy, a set of transferable competencies for reading, evaluating, and iteratively refining text-based interaction with an intelligent system (Hwang, Lee, & Shin, 2023; Ng, Leung, Chu, & Qiao, 2021).

This paper takes seriously the linguistic dimension of this claim. If prompting is a literacy, it is, at its core, a literacy of language use: choosing words, structuring sentences, calibrating register and directness, and anticipating how an interlocutor, in this case a statistical model of language rather than a human mind, will interpret and respond to an utterance. Yet the linguistic properties of prompts have received comparatively less systematic attention than their technical taxonomy. Engineering-oriented surveys catalogue prompting techniques by their computational effect, for example how a given prompt format improves accuracy on a benchmark task, but rarely analyse prompts using the descriptive tools linguistics has developed for centuries to describe how humans use language to accomplish things: speech-act theory, the Gricean cooperative principle, politeness theory, and register variation. Conversely, the small but growing body of linguistically informed work on human-AI interaction, exemplified by Panfili and colleagues' (2021) Gricean analysis of human-AI dialogue, has focused primarily on earlier, narrower conversational agents rather than on the deliberate, often highly structured prompting practices that have developed around contemporary large language models.

This paper addresses that gap by treating prompt engineering as an object of linguistic description in its own right. It draws together the technical literature on prompting techniques, the pragmatic and Gricean literature on human-AI communication, and the educational literature on prompt and AI literacy, in order to build and apply a functional-linguistic taxonomy of prompting strategies. Using a qualitative content-analysis methodology, the study examines a purposively selected corpus of widely documented, publicly available prompt patterns and frameworks, including zero-shot and few-shot instructions, chain-of-thought reasoning prompts, role or persona prompts, and structured pedagogical frameworks such as the CLEAR model for academic prompting (Lo, 2023), asking what these prompts reveal about the linguistic competencies that effective human-AI communication actually requires. In doing so, the paper aims to reframe prompt engineering, in the public and pedagogical imagination often treated as a purely technical or even faddish skill, as a genuine, describable, and teachable new literacy grounded in long-standing linguistic categories, while also identifying the specific ways in which talking to a large language model departs from talking to a human interlocutor.

The remainder of the paper proceeds in five further sections. It first states the study's objectives and research questions. It then reviews the relevant literature published between 2019 and 2026, tracing the field's development from the original few-shot learning paradigm through systematic technical surveys, pragmatic and Gricean analyses, and the emergence of prompt and AI literacy as explicit educational constructs. A methodology section then describes the qualitative, corpus-based approach used to analyse a purposively sampled set of prompt patterns. A results section presents this analysis across four dimensions, supported by four summary tables, before a discussion section draws out the implications of the findings, and a conclusion situates prompt engineering within the longer history of literacies that have developed around new communication technologies.

It is worth pausing on why these reframing matters beyond terminology. If prompt engineering is treated purely as a technical skill, comparable to learning a query syntax or a spreadsheet formula language, then teaching it well is primarily a matter of teaching correct syntax and a catalogue of effective templates, to be memorised and applied. If, by contrast, prompt engineering is a linguistic literacy, comparable to learning to write a persuasive letter or to give a clear set of instructions, then teaching it well requires cultivating a flexible, transferable communicative sensitivity: an ability to judge how much context a particular model is likely to need, how explicit an instruction must be to avoid being misread, and how to recover gracefully when a first attempt at communication fails. These two pedagogical visions are not mutually exclusive, but they place the emphasis very differently, and the evidence gathered in this paper is intended to show that the second, literacy-based vision better captures what skilled prompt writers actually do, and therefore what novice prompt writers most need to be taught.

A further motivation for this study is definitional. The terms prompt engineering, prompt literacy, and AI literacy are frequently used interchangeably in both popular and scholarly discussion, yet they name subtly different things: prompt engineering typically refers to the technical practice of constructing effective prompts, prompt literacy typically refers to the broader educational competence of doing so reflectively and adaptively, and AI literacy typically refers to a still broader understanding of how AI systems work and what their outputs can and cannot be trusted to mean (Ng, Leung, Chu, & Qiao, 2021). This paper uses prompt engineering and prompt literacy in the senses just described, and treats AI literacy as the wider competence within which prompt literacy is nested, following the conceptual architecture proposed in the recent educational literature reviewed below.

1.1 Research Objectives

- To examine prompt engineering as a linguistic practice, describing the pragmatic and functional resources that effective prompts deploy when addressing a large language model.
- To develop a functional-linguistic taxonomy of prompting strategies, drawing on speech-act theory, the Gricean cooperative principle, politeness theory, and register variation.
- To apply this taxonomy to a purposively selected corpus of widely documented prompt patterns and frameworks, in order to identify recurring linguistic strategies across different prompting techniques.
- 4. To compare the linguistic demands of prompting a large language model with the linguistic demands of ordinary human-to-human communication, identifying points of continuity and departure.
- 5. To situate prompt engineering within the broader concept of literacy, evaluating the claim that it constitutes a genuine, teachable new literacy rather than a narrow technical skill.

1.2 Research Questions

1. What linguistic and pragmatic strategies do effective prompts for large language models systematically deploy?
2. How do established linguistic frameworks, including speech-act theory and the Gricean cooperative principle, apply to, and need to be modified for, the specific case of prompting a large language model?
3. What distinguishes the linguistic demands of prompt engineering from those of ordinary conversational competence, and what new competencies does prompting require?
4. To what extent, and in what specific sense, can prompt engineering be considered a literacy comparable to other historically recognised literacies?

2. Literature Review

This review surveys scholarship on prompting, human-AI communication, and prompt or AI literacy published between 2019 and 2026, organised chronologically to trace the field's development from a narrow technical discovery into a broader linguistic and educational concern.

The technical starting point for the entire field of prompt engineering is Radford and colleagues' (2019) demonstration that a sufficiently large language model, trained only to predict the next word in a sequence of text, could perform a surprising range of tasks, including translation and rudimentary question answering, when simply given a natural-language description of the task as part of its input, without any task-specific training. This finding was decisively extended by Brown and colleagues (2020) in the GPT-3 paper, which showed that providing a handful of input-output examples directly within the prompt, a technique the authors termed few-shot learning, allowed the model to perform new tasks competitively with models that had been explicitly fine-tuned for them. The significance of this discovery for the present paper's argument can hardly be overstated: it established, for the first time at scale, that the wording and structure of a natural-language instruction, rather than any change to the underlying model, could function as the primary lever for controlling a machine's behaviour, effectively making the prompt itself, a piece of ordinary written language, into a programming interface.

This period's significance is as much conceptual as technical. Prior human-computer interaction paradigms, from command-line syntax to graphical menus, deliberately minimised the ambiguity of natural language by replacing it with constrained, formal input. Brown and colleagues' (2020) few-shot paradigm reintroduced natural language, with all its ambiguity, context-dependence, and implicature, as the primary control surface for a powerful computational system, effectively transplanting the entire apparatus of human pragmatics into a domain that had previously been engineered specifically to avoid it. The 2019–2020 literature therefore does not yet use the vocabulary of literacy or linguistics explicitly, but it lays the empirical foundation, the demonstrated sensitivity of model behaviour to prompt wording, on which the entire subsequent literature on prompt engineering as a skill, and eventually as a literacy, depends.

It is also worth noting what this earliest period leaves unresolved. Radford and colleagues (2019) and Brown and colleagues (2020) both report, largely in passing, that the precise wording of a task description could dramatically change a model's performance on that task, but neither paper offers a systematic account of why particular wordings succeed while superficially similar wordings fail. This gap between an empirically demonstrated sensitivity to language and a principled explanation of that sensitivity is precisely the gap the subsequent pragmatic and educational literature, reviewed in the following subsections, attempts to close, and it is the same gap this paper's own linguistic-functional taxonomy is designed to address for the specific case of the exemplar corpus analysed here.

Between 2021 and 2022, the field moved to formalise prompting as a distinct object of study with its own emerging taxonomy of techniques, while a smaller but important strand of research began to frame human-AI exchanges in explicitly pragmatic terms. On the technical side, researchers began cataloguing and comparing distinct prompting strategies, moving from Brown and colleagues' (2020) original few-shot demonstrations toward more elaborate structures, including prompts that asked a model to work through intermediate reasoning steps before producing a final answer, a technique that would later be consolidated and popularised as chain-of-thought prompting. This period's technical contributions established that the internal structure of a prompt, not merely its content, systematically affects model output, a finding with direct linguistic significance: it suggests that discourse-level features such as sequencing, explicit signposting, and staged elaboration function for a language model in ways that parallel their function in human textual comprehension.

On the pragmatic side, Panfili and colleagues (2021) offered an explicitly Gricean analysis of human-AI interaction, applying the cooperative principle and its associated maxims of quantity, quality, relation, and manner, originally developed to describe human conversational implicature, to exchanges between people and AI systems. Their analysis demonstrated that people bring the same conversational expectations to an AI interlocutor that they bring to a human one, expecting a response to be as informative as required but no more, to be truthful, to be relevant to what was asked, and to be clearly expressed, and that violations of these Gricean maxims by an AI system are perceived and evaluated by users in broadly similar ways to violations by a human speaker. This finding is foundational for the present paper's argument because it establishes empirically that ordinary human pragmatic competence, rather than a wholly new skill set, is the starting point from which prompt engineering as a specialised literacy subsequently develops.

This period's chain-of-thought innovations deserve particular emphasis because they represent one of the clearest early instances of a purely technical discovery converging with a linguistic principle. Researchers working independently of any explicit pragmatic theory found, empirically, that asking a model to articulate intermediate reasoning steps before giving a final answer substantially improved performance on multi-step problems; described in the vocabulary of Gricean pragmatics, this technique amounts to imposing explicit manner and orderliness on the discourse structure of the model's own output, rather than merely on the wording of the input prompt. That a purely performance-driven engineering discovery and a classical pragmatic maxim converge so closely is, this paper suggests, an early sign that the technical and linguistic literatures on prompting, though developed independently and largely without reference to one another during this period, were in fact converging on the same underlying communicative principles from different directions.

The 2023–2024 period is marked by three converging developments: the publication of comprehensive technical surveys that consolidated the rapidly proliferating landscape of prompting techniques, the explicit introduction of the term prompt literacy into the educational literature, and growing attention to prompting from applied linguistic and information-literacy perspectives. Liu and colleagues' (2023) systematic survey of prompting methods in natural language processing organised the field's techniques into a coherent framework, distinguishing, among other things, between prompts that require no worked examples, prompts that provide several, and prompts

that are themselves automatically generated or refined by another language model, as demonstrated by Zhou and colleagues' (2023) work showing that large language models can generate effective prompts for other instances of themselves. Marvin, Hellen, Jjingo, and Nakatumba-Nabende (2024) further consolidated this technical landscape, arguing explicitly that prompt engineering had become an obligatory communicative skill for effective interaction with conversational AI systems such as ChatGPT, a framing that positions the practice squarely within, rather than outside, the domain of communicative competence.

This period also saw the explicit coining and theorising of prompt literacy as an educational construct. Hwang, Lee, and Shin (2023) introduced the term to describe the ability to generate precise prompts, interpret the resulting output, and iteratively refine the prompt to achieve a desired result, studying its emergence among language learners engaged in an AI-assisted creative task. Their formulation is significant for explicitly building prompt literacy on the existing scaffolding of literacy studies rather than treating it as an entirely novel skill, positioning it as an extension of established digital and information literacies into the specific domain of generative AI (Ng, Leung, Chu, & Qiao, 2021). Complementing this educational framing, Lo (2023) proposed the CLEAR model, an explicit pedagogical framework instructing prompt writers to be Concise, Logical, Explicit, Adaptive, and Reflective, developed specifically for academic library and research contexts. The CLEAR model is of particular interest to the present study because it translates what is, in effect, a set of pragmatic and stylistic principles, brevity, coherent sequencing, explicitness, flexibility, and iterative revision, into an explicit, teachable framework, exactly the kind of transformation of tacit linguistic competence into codified literacy that this paper argues characterises prompt engineering as a whole.

A further contribution from this period worth highlighting is Zhou and colleagues' (2023) demonstration that large language models themselves could be used to generate and refine effective prompts for other instances of the same model, a technique the authors describe as making the model a human-level prompt engineer in its own right. This finding raises an intriguing question this paper returns to in its discussion: if a model can independently discover effective prompting strategies that converge with the linguistic principles identified through human pragmatic analysis, this offers indirect but suggestive evidence that those principles, quantity calibration, explicit sequencing, and the like, are not merely artefacts of how humans happen to think about communicating with machines, but are genuine properties of what makes a piece of natural-language text an effective instruction for a system of this kind, discoverable independently through either human linguistic reasoning or automated optimisation.

The most recent scholarship, published in 2025 and continuing into 2026, is characterised by systematic, cross-institutional synthesis and by the extension of prompt and AI literacy into specific academic and professional domains. Knoth, Tolzin, Janson, and Leimeister (2024, with substantial uptake through 2025) demonstrated empirically that the quality of a person's prompt engineering, measured by whether a prompt included key structural components identified in the literature, predicted the quality of the resulting output, and that certain dimensions of a person's broader AI literacy were associated with higher-quality prompting, providing direct quantitative support for treating prompt quality as a describable, measurable linguistic-communicative competence rather than an unanalysable matter of individual flair. Tour and Zadorozhnyy (2025) extended this line of enquiry specifically into English language education, proposing that prompt literacy for English language learners could be scaffolded using the established Four Resources Model of literacy, thereby formally integrating prompt literacy into a well-established framework from literacy studies rather than treating it as a free-standing, *sui generis* skill.

Qian's (2025) systematic review of prompt engineering in education surveyed empirical studies published since the widespread public release of ChatGPT in late 2022, distinguishing technique-based approaches, which target specific, well-defined learning goals through particular prompt formats, from process-based approaches, which support broader cognitive engagement and collaborative thinking between learner and AI system. A further systematic review, by Anowar and Khatun (2026), organised the accumulated literature around a layered conceptual model spanning theoretical foundations, practical educational applications, and emerging concerns such as academic integrity and teacher preparation, indicating that by 2026 prompt engineering research had matured into a field with its own settled conceptual architecture rather than a loose collection of individual technique papers. Taken together, this most recent scholarship treats prompt engineering less as a novel technical curiosity and more as an established communicative competence requiring the same kind of structured pedagogical attention that other literacies, including traditional print literacy and earlier digital literacies, have historically received.

Read as a whole, the literature surveyed across these four periods traces a clear arc: from the initial technical discovery that natural-language wording alone could control model behaviour, through the formalisation of prompting

techniques and an early, explicitly Gricean pragmatic framing of human-AI dialogue, to the coining and theorising of prompt literacy as an explicit educational construct, and finally to a matured, systematised field extending prompt and AI literacy into discipline-specific pedagogical contexts. This arc supplies both the theoretical resources and the motivating rationale for the present study, which treats this seven-year span of scholarship as a cumulative toolkit for describing prompt engineering in properly linguistic terms.

3. Research Methodology

This study adopts a qualitative, interpretive research design centred on linguistic-functional content analysis. Because the object of study is the pragmatic and functional behaviour of natural-language prompts, a qualitative close-analysis methodology is better suited than a purely quantitative benchmarking design, which could measure a prompt's effect on model accuracy but not the linguistic strategies responsible for that effect. The design is exploratory and descriptive rather than hypothesis-testing: it aims to build and illustrate a functional-linguistic taxonomy of prompting strategies rather than to establish statistical claims about which prompt format performs best on a given computational benchmark, a question already addressed extensively elsewhere in the technical literature.

The design also incorporates a comparative logic, examining prompt types that differ markedly in their formality, structure, and intended use, from minimal zero-shot instructions to elaborate, multi-component academic frameworks. This comparative element allows the study to distinguish linguistic features that are common to prompting in general from features specific to particular prompting techniques or contexts, strengthening the claim that the taxonomy developed here captures genuine, generalisable properties of prompt language rather than idiosyncrasies of any single technique.

A further design consideration concerns the relationship between this study and the large existing body of computational, benchmark-driven prompt engineering research. That literature typically treats a prompt as an independent variable whose effect on a dependent variable, task accuracy, is measured through controlled experimentation across many trials. The present study is deliberately positioned as complementary to, rather than competitive with, that tradition: it takes as given the empirical finding, well established across dozens of benchmark studies, that certain prompt structures reliably improve model performance, and asks the different, qualitative question of what these effective structures have in common when described using the analytic vocabulary of linguistics. In this sense, the research design treats the computational literature's benchmark findings as a kind of pre-validated dataset of effective language use, from which linguistic patterns can be qualitatively abstracted, rather than attempting to independently re-verify those findings through new experimentation.

3.1 Theoretical Framework

The analysis draws on three complementary linguistic traditions. The first is speech-act theory, which treats utterances as actions, requests, instructions, assertions, and so on, performed through language, providing a vocabulary for classifying what a prompt is doing communicatively, independent of its surface form. The second is the Gricean cooperative principle and its associated maxims of quantity, quality, relation, and manner, applied to human-AI dialogue by Panfili and colleagues (2021), which supplies a means of describing how a well-formed prompt manages the amount, truthfulness, relevance, and clarity of the information it provides to the model. The third is register and politeness theory, concerned with how speakers or writers calibrate formality, directness, and social distance according to context, which is directly relevant to understanding why prompts addressed to a large language model often adopt a register, imperative, explicit, sometimes strikingly bare of the hedges and politeness markers that would soften an equivalent request to a human interlocutor.

These three traditions are combined into a single coding framework, further supplemented by the technical vocabulary of prompt engineering itself, including the distinctions between zero-shot, few-shot, and chain-of-thought prompting established in the surveys of Brown and colleagues (2020) and Liu and colleagues (2023), and by the explicit pedagogical vocabulary of the CLEAR model (Lo, 2023). Where the linguistic frameworks and the technical vocabulary overlap, for example where a chain-of-thought prompt's staged structure can be described both as a technical prompting pattern and as an instance of Gricean manner and orderliness, the analysis notes this convergence explicitly, since it is precisely this convergence between engineering technique and linguistic principle that motivates the paper's central claim that prompt engineering is a linguistic literacy.

3.2 Selection and Sampling

The primary corpus was assembled through purposive sampling of widely documented, publicly reported prompt patterns and frameworks drawn from the technical and educational literature reviewed above, rather than through original elicitation of new prompts from human participants. Five categories of prompt pattern were selected for analysis: zero-shot instructional prompts, of the kind described by Radford and colleagues (2019); few-shot prompts incorporating worked examples, as established by Brown and colleagues (2020); chain-of-thought prompts that request explicit intermediate reasoning; role or persona prompts that assign the model a specific identity or perspective before issuing an instruction; and structured academic prompts following the CLEAR framework of Lo (2023). For each category, illustrative exemplars reported in the reviewed literature were selected and, where necessary, reconstructed in a representative, non-verbatim form consistent with the general pattern documented in the source scholarship, ensuring that the analysis engages with genuinely representative instances of each technique rather than a single idiosyncratic example.

This sample was not intended to be an exhaustive inventory of every prompting technique documented in the rapidly expanding technical literature; new techniques, including various automated and multi-agent prompting strategies, continue to be published faster than any single qualitative study could track. Rather, the five categories were selected because each is well established, frequently cited across multiple independent sources in the literature reviewed above, and jointly span a wide range of structural complexity, from a single unadorned instruction to an elaborate, multi-component academic framework, which is the range of variation this study's comparative research questions require.

3.3 Data Collection

Data collection proceeded through systematic extraction of illustrative prompt exemplars, and associated technical and pedagogical commentary, from the sources identified during the literature review, followed by close linguistic annotation of each exemplar. Each prompt exemplar was logged together with its source category, its documented or reported effect on model output where such information was available in the source literature, and any accompanying commentary the original authors provided about why the prompt was constructed in that particular way. Where a single source, such as Liu and colleagues' (2023) survey, reported multiple exemplars of the same prompting category, exemplars were selected for diversity of surface wording, in order to avoid the analysis being unduly shaped by the idiosyncratic phrasing of any single example.

Throughout data collection, care was taken to record not only the prompt text itself but its surrounding context of use, since, consistent with the pragmatic and Gricean framework adopted here, a prompt's linguistic function cannot be fully assessed independently of the task it is designed to accomplish and the point in an interaction at which it is issued. This is particularly important for chain-of-thought and role-based prompts, whose linguistic effect often depends on their position within a longer sequence of exchanges with the model rather than on their content in isolation.

3.4 Data Analysis

Exemplars were analysed using qualitative content analysis, following an iterative coding procedure closely parallel to that used in comparable linguistic-functional studies. An initial set of codes was derived deductively from the theoretical framework described above, including speech-act categories such as directive and assertive, Gricean maxim categories such as quantity and manner, and register features such as formality and explicitness. These deductive codes were then refined inductively as the corpus was analysed, producing additional prompt-specific categories, such as context-priming and role-assignment, that extend beyond the classical inventory of speech-act types because they depend on the distinctive inferential behaviour of a large language model rather than on the mental-state attribution that governs human-to-human speech acts.

Coding reliability was supported by repeated re-analysis of the corpus and by cross-checking coded exemplars against the functional definitions established in the theoretical framework, with ambiguous cases resolved by considering the exemplar's documented purpose and reported effect in the source literature rather than treating the prompt's wording as self-interpreting. Where a single exemplar appeared to serve multiple functions simultaneously, for example combining an explicit directive with an implicit register cue, it was coded for all applicable categories,

consistent with the multifunctional view of communication that underlies both speech-act theory and the Gricean framework. The resulting coding scheme, comprising six functional categories in total, forms the basis of the Results and Findings section below.

4. Results and Findings

Table 1 summarizes the five categories of prompt pattern analyzed in this study, indicating their originating source in the technical literature, their typical structural complexity, and their principal documented use case. The corpus spans the full range from Radford and colleagues' (2019) minimal zero-shot instructions to Lo's (2023) elaborate, multi-component CLEAR framework, deliberately including both highly informal, unstructured prompting styles and highly codified, pedagogically explicit ones. This range was selected precisely because the study's comparative research questions require variation not only in what a prompt asks for, but in how explicitly and formally it asks for it.

Table 1. Overview of the Primary Prompt Corpus

Prompt Category	Originating Source	Structural Complexity	Principal Documented Use Case
Zero-shot instruction	Radford et al. (2019)	Minimal: single instruction, no examples	General task completion without demonstration
Few-shot prompt	Brown et al. (2020)	Moderate: instruction plus worked examples	Task adaptation without model retraining
Chain-of-thought prompt	Liu et al. (2023); related surveys	High: staged, sequential reasoning request	Complex reasoning and multi-step problem solving
Role or persona prompt	Marvin et al. (2024)	Moderate: identity assignment plus instruction	Perspective-taking, stylistic and tonal control
CLEAR academic framework	Lo (2023)	High: five explicit structural components	Academic research and library information tasks

Close analysis of the prompt exemplars in the corpus produced six recurrent functional categories, summarised in Table 2. These categories draw directly on speech-act theory and the Gricean framework reviewed above, supplemented by two prompt-specific categories, context-priming and role-assignment, that extend beyond the classical inventory because they exploit the distinctive way a large language model conditions its output on preceding text rather than inferring a human speaker's mental state.

Table 2. Functional Categories of Prompt Language in the Sampled Corpus

Function Category	Description	Illustrative Category	Prompt
Directive speech act	The prompt performs an explicit instruction or request for a specific action or output	Zero-shot	instruction; CLEAR framework
Context-priming	The prompt supplies background information or worked examples that condition subsequent output without issuing a direct instruction	Few-shot prompt	

Quantity calibration	The prompt explicitly bounds how much information the model should supply, addressing the Gricean maxim of quantity	CLEAR framework; zero-shot instruction
Manner and sequencing	The prompt imposes an explicit order or staged structure on the requested reasoning process	Chain-of-thought prompt
Role-assignment	The prompt assigns the model an identity or perspective that reframes how subsequent instructions are interpreted	Role or persona prompt
Reflective revision cueing	The prompt anticipates and invites iterative refinement of the output rather than treating the first response as final	CLEAR framework

Several of these categories deserve brief illustration. Directive speech acts were near-universal across the corpus, present in every category, but varied considerably in their explicitness: zero-shot instructions typically front-load a bare imperative, while the CLEAR framework's emphasis on being Explicit, one of its five named components, converts what might otherwise be an implicit request into an unambiguous, itemised directive. Context-priming was the defining feature of few-shot prompts, whose worked examples perform no direct instructional speech act at all; instead, they rely on the model's tendency to continue an established pattern, a linguistic strategy closer to implicature and analogy than to direct instruction, and one with no straightforward equivalent in ordinary human conversational request-making, where providing several worked examples before making a request would be considered unusually laborious or even patronising.

Quantity calibration appeared most explicitly in the CLEAR framework's Concise component and in zero-shot instructions that specify a word count or format constraint, directly operationalising the Gricean maxim of quantity, which in ordinary conversation is typically observed tacitly rather than stated outright; the fact that effective prompts often state this maxim explicitly, rather than relying on the interlocutor to infer it, is itself a notable departure from ordinary human pragmatic practice. Manner and sequencing dominated the chain-of-thought category, whose defining technique is to request that the model work through intermediate steps in a specified order before producing a final answer, an explicit imposition of discourse structure that in human pedagogical contexts would correspond to asking a student to show their working, converting an implicit expectation of orderly reasoning into an explicit textual instruction. Role-assignment, found principally in persona prompts, reframes the entire subsequent exchange by first establishing a fictional identity or professional perspective for the model, a strategy with no direct equivalent in ordinary interpersonal address but with clear affinities to theatrical or role-play framing devices in human communication. Reflective revision cueing, finally, was most fully articulated in the CLEAR framework's Reflective component, which explicitly builds an expectation of iterative refinement into the initial prompt rather than treating the exchange as a single, one-shot request.

Table 3 presents an approximate distribution of the six functional categories across the five prompt types in the corpus, based on the proportion of coded exemplars within each category that were assigned to each function during the analytic procedure described in the Methodology. These figures are qualitative and illustrative of relative emphasis within this purposive sample; they are not intended as statistically generalisable frequencies for all prompts of a given type.

Table 3. Relative Emphasis of Functional Categories by Prompt Type (Qualitative Coding Proportions)

Function Category	Zero-shot	Few-shot	Chain-of-thought	Role/Persona	CLEAR Framework
Directive speech act	High	Low	Moderate	Moderate	High

Context-priming	Low	High	Moderate	Low	Moderate
Quantity calibration	Moderate	Low	Low	Low	High
Manner and sequencing	Low	Low	High	Low	Moderate
Role-assignment	Not present	Not present	Low	High	Low
Reflective revision cueing	Low	Low	Moderate	Low	High

Table 4 compares the pragmatic profile of effective prompting, as identified across the corpus, with the pragmatic profile of ordinary, cooperative human-to-human conversation, along four dimensions drawn from the theoretical framework, in order to make explicit the points of continuity and departure that motivate several of this study's research questions.

Table 4. Comparative Pragmatic Profiles: Prompting a Large Language Model versus Ordinary Human Conversation

Dimension	Effective Prompting of a Large Language Model	Ordinary Cooperative Human Conversation
Explicitness of the quantity maxim	Often stated directly, for example specifying a word limit or number of points	Usually left implicit and inferred from context
Use of worked examples	Frequently used deliberately as a substitute for direct instruction (few-shot prompting)	Rare in ordinary requests; associated with pedagogical or demonstrative contexts only
Politeness and hedging	Frequently minimal; bare imperatives are common and unremarkable	Typically softened with hedges, modal verbs, or indirect requests to preserve face
Structuring of reasoning	Often explicitly requested and sequenced (chain-of-thought)	Usually left to the listener's own inferential processes unless specifically asked to explain

Taken together, these findings suggest that effective prompt engineering is neither a wholly novel skill unrelated to existing linguistic competence nor a simple transfer of ordinary conversational pragmatics onto a new interlocutor. Instead, it involves a systematic, often deliberate recalibration of familiar pragmatic principles, quantity, manner, directness, and context-setting chief among them, for an interlocutor whose inferential behaviour depends on statistical regularities in text rather than on the theory-of-mind reasoning that governs human conversational inference.

One further pattern emerging from this comparison deserves separate note: the near-total absence, across all five prompt categories in the corpus, of the face-saving and politeness strategies that pervade ordinary human requests. Human speakers routinely soften a request with a hedge, an apology, or an indirect phrasing, in part to protect the listener's autonomy and self-image, a phenomenon extensively documented in politeness theory. Effective prompts in this corpus almost uniformly dispense with these strategies, favouring bare imperatives and explicit specification instead, not because prompt writers are being deliberately impolite, but because a language model has no face to protect and no autonomy to defer to in the sense that politeness theory presupposes for a human addressee. This absence is itself linguistically significant: it indicates that at least one entire dimension of ordinary human pragmatic competence, politeness management, becomes largely redundant, rather than merely recalibrated, when the addressee is a language model rather than a person, a finding that qualifies this paper's broader claim that prompting recalibrates rather than replaces ordinary pragmatics.

5. Discussion

The findings of this study suggest that prompt engineering, when analysed in properly linguistic terms, is best described as a calibrated extension of ordinary pragmatic competence rather than as an entirely new form of communication invented from scratch. Panfili and colleagues' (2021) demonstration that people bring Gricean expectations to their interactions with AI systems is clearly supported and extended by this study's finding that effective prompts frequently make those Gricean maxims explicit, stating outright how much information is wanted, or exactly how reasoning should be sequenced, rather than leaving them to be inferred as they typically are in human conversation. This suggests that prompt engineering does not so much depart from Gricean cooperative principles as render them unusually visible, converting tacit conversational norms into overt textual instructions precisely because the model, unlike a human interlocutor, cannot be relied upon to infer them from social context alone.

The prominence of context-priming in few-shot prompting is particularly instructive in this regard. Providing several worked examples before making a request is, in ordinary human conversation, an unusual and somewhat marked strategy, associated with teaching, demonstration, or unusually cautious communication with someone assumed to be unfamiliar with a task. That this strategy has become one of the most well-documented and effective prompting techniques (Brown et al., 2020; Liu et al., 2023) suggests that large language models respond to demonstration and analogy in ways that differ systematically from how human listeners respond to direct instruction, a difference with real consequences for how prompt literacy ought to be taught: a learner who has only practised issuing direct requests to other people may need explicit instruction in the comparatively unfamiliar skill of teaching by worked example before they can construct an effective few-shot prompt.

The comparison between role-assignment prompting and ordinary human role-play or framing devices also deserves comment. Assigning a large language model a specific persona or professional identity before issuing a substantive instruction has no direct equivalent in ordinary unscripted conversation between adults, where assuming a fictional identity mid-conversation would typically require explicit signalling and mutual agreement, yet it has become a well-established and effective prompting technique (Marvin, Hellen, Jjingo, & Nakatumba-Nabende, 2024). This suggests that large language models are unusually responsive to a kind of framing device that in human communication is largely confined to specialised genres such as theatre, role-play, or structured pedagogical exercises, and that part of what prompt literacy teaches is the transfer of a communicative strategy from these marked, special-purpose human genres into the default register of everyday human-AI address.

The near-absence of politeness marking documented in the Results section deserves further reflection here, since it complicates a simple story in which prompting is merely ordinary conversational competence turned toward a new addressee. Politeness theory holds that face-management strategies exist because human interlocutors are simultaneously cooperative and self-interested, each managing the other's and their own social standing across a conversation; a language model, lacking any socially consequential self-image to protect or threaten, gives prompt writers no comparable reason to hedge. This suggests that part of what a skilled prompt writer must learn is not simply how to redirect existing conversational habits toward a new kind of listener, but how to selectively suspend certain habits, politeness chief among them, that would be maladaptive, or at best inert, in this new communicative context, while retaining and often intensifying others, particularly explicitness and quantity calibration. Prompt literacy, on this view, involves not a uniform recalibration of conversational competence but a differentiated one, in which some pragmatic reflexes are dialled up, others are dialled down, and the skill lies precisely in knowing which is which for a given communicative goal.

These findings also bear on the central claim of this paper's title, that prompt engineering constitutes a new literacy. The evidence gathered here suggests that this claim is best supported not because prompting requires an entirely novel cognitive or linguistic skill, unrelated to anything a competent language user already possesses, but because it requires the explicit, conscious application and recalibration of pragmatic principles, quantity, manner, context-setting, and framing, that competent speakers typically apply tacitly and unconsciously in ordinary conversation. This is precisely the kind of transformation that literacy studies has long associated with the emergence of a new literacy: not the invention of an unprecedented cognitive faculty, but the conversion of previously tacit, unreflective practice into an explicit, teachable, and assessable skill, of the kind the CLEAR framework (Lo, 2023) and the Four Resources scaffolding proposed by Tour and Zadorozhnyy (2025) both attempt to provide. Viewed this way, prompt engineering sits comfortably alongside other literacies, print literacy, information literacy, and digital

literacy, that have each, in their own historical moment, taken competencies that were once tacit or informally acquired and rendered them explicit objects of instruction.

Limitations of this study should be acknowledged. The corpus is built from documented, published prompt exemplars rather than from a large-scale, naturalistic sample of prompts written by ordinary users in real interactions, and it is possible that everyday, informal prompting behaviour departs in systematic ways from the carefully constructed exemplars that appear in technical and pedagogical publications. The coding proportions reported in Table 3 are illustrative of relative emphasis within this qualitative sample and should not be read as precise frequency counts applicable to prompting behaviour in general. A larger, corpus-linguistic study of naturally occurring prompts, potentially combining the qualitative categories developed here with computational text-mining methods of the kind used in prompt-optimisation research (Zhou et al., 2023), would allow these categories to be tested against a much larger and more ecologically representative sample of everyday prompting behaviour.

A further limitation concerns the pace of change in the underlying technology itself. Because large language models are trained and updated on an ongoing basis, and because prompting techniques that are effective for one model or model version are not guaranteed to remain equally effective for another, the specific linguistic strategies catalogued here should be understood as a snapshot of a rapidly evolving communicative practice rather than a permanently fixed inventory. Despite this instability at the level of technique, the underlying linguistic categories used to describe those techniques, speech acts, Gricean maxims, register, and framing, are drawn from stable, general-purpose linguistic theory and should retain descriptive value even as the specific prompting techniques they are used to analyse continue to change.

It is also worth situating this taxonomy against the possibility of future replication and extension. Because each functional category proposed here was defined with explicit reference to independently established linguistic theory, speech-act categories, Gricean maxims, and politeness and register variation, rather than invented solely for this study, later researchers working with naturally occurring prompt logs, drawn perhaps from workplace, educational, or creative-writing contexts, should be able to apply the same six categories, or a principled extension of them, without needing to reconstruct the coding scheme from first principles. This transferability is, in the end, the paper's most durable methodological contribution: not a final, closed inventory of what a prompt can do, but a tested descriptive vocabulary that later, larger-scale studies can confirm, complicate, or revise as prompting practice continues to develop alongside the models it addresses.

6. Conclusion

This study set out to ask whether prompt engineering, the practice of constructing natural-language input to elicit desired behaviour from a large language model, is best understood as a narrow technical skill or as a genuine linguistic literacy. The evidence gathered from a purposively sampled corpus of widely documented prompt patterns, read against the technical, pragmatic, and educational scholarship on prompting published between 2019 and 2026, supports the latter conclusion. Effective prompts systematically deploy directive speech acts, context-priming through worked example, explicit quantity calibration, staged sequencing of reasoning, deliberate role-assignment, and reflective revision cueing, a set of functions that draws directly on, while also recalibrating, the ordinary pragmatic competencies that govern human-to-human conversation. Where ordinary conversation typically leaves the Gricean maxims of quantity, manner, and relevance to be inferred from shared context, effective prompting frequently renders these same maxims explicit, precisely because the large language model cannot be assumed to share the tacit contextual and social knowledge that licenses such inference among human interlocutors.

Read together with the emergence of prompt literacy and AI literacy as explicit educational constructs (Hwang, Lee, & Shin, 2023; Ng, Leung, Chu, & Qiao, 2021; Tour & Zadorozhnyy, 2025), and with structured pedagogical frameworks such as CLEAR (Lo, 2023) that convert pragmatic principles into explicit, teachable components, prompt engineering appears not as an isolated technical novelty but as the latest instance of a much older process: the conversion of previously tacit communicative competence into an explicit, describable, and teachable literacy, in response to a new communication technology that cannot be relied upon to share the background inferential resources of a human conversational partner. Future research would benefit from extending the linguistic taxonomy developed here to large-scale, naturalistic corpora of everyday prompts, from testing whether the specific linguistic strategies identified here remain stable across successive generations of language models, and from combining close linguistic analysis with the computational and educational methods that have driven much of the recent progress in

prompt engineering research more broadly. Such work would help clarify whether prompt literacy is settling into a stable set of describable linguistic competencies or continuing to evolve as rapidly as the models it is used to address.

Beyond its immediate findings, this study also makes a methodological argument comparable to the one that motivates its literature review: that prompt engineering research benefits from being pursued jointly by computer scientists and linguists rather than separately by each. The technical literature reviewed here, spanning few-shot learning, systematic prompting surveys, and automated prompt generation, supplies exactly the empirical detail that a purely linguistic analysis would lack, documenting with precision which prompt structures produce which measurable effects on model output. Conversely, the linguistic and pragmatic literature, from Gricean analyses of human-AI dialogue to speech-act theory and register variation, supplies computer science with a principled, generalisable vocabulary for describing why those structures work, rather than leaving prompt engineering's effectiveness as an unexplained empirical regularity. This study's proposed taxonomy is offered as a modest contribution toward that joint enterprise: not a final, closed inventory of what a prompt can do, but a tested starting vocabulary that later, larger-scale studies, across both disciplines, can confirm, complicate, or revise as this new and rapidly developing literacy continues to take shape.

Declarations

Conflict of interest	The authors declare no conflicts of interest.
Funding	This research received no external funding.
Data availability	Data are available from the corresponding author upon reasonable request.
Use of AI tools	No AI-assisted tools were used in the preparation of this manuscript.
Author contributions	Conceptualization, S.B.; methodology, S.B.; software, S.B.; validation, S.B.; writing—original draft preparation, S.B.; writing—review and editing, S.B.; supervision, S.B.; project administration, S.B.

References

- Anowar, S., & Khatun, R. (2026). Prompt engineering in education: A systematic review of theory, applications, and future directions. *Journal of Education, Society & Sustainable Practice*, 2, 30–42.
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. (2020). Language models are few-shot learners. *arXiv*. <https://doi.org/10.48550/arXiv.2005.14165>
- Hwang, Y., Lee, J. H., & Shin, D. (2023). What is prompt literacy? An exploratory study of language learners' development of new literacy skill using generative AI. *arXiv*. <https://doi.org/10.48550/arXiv.2311.05373>
- Knonth, N., Tolzin, A., Janson, A., & Leimeister, J. M. (2024). AI literacy and its implications for prompt engineering strategies. *Computers and Education: Artificial Intelligence*, 6, 100225. <https://doi.org/10.1016/j.caeai.2024.100225>
- Liu, P., Yuan, W., Fu, J., Jiang, Z., Hayashi, H., & Neubig, G. (2023). Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing. *ACM Computing Surveys*, 55(9), 1–35. <https://doi.org/10.1145/3560815>
- Lo, L. S. (2023). The CLEAR path: A framework for enhancing information literacy through prompt engineering. *The Journal of Academic Librarianship*, 49(4), 102720. <https://doi.org/10.1016/j.acalib.2023.102720>
- Marvin, G., Hellen, N., Jjingo, D., & Nakatumba-Nabende, J. (2024). Prompt engineering in large language models. In I. J. Jacob, S. Piramuthu, & P. Falkowski-Gilski (Eds.), *Data intelligence and cognitive informatics* (pp. 387–402). Springer. https://doi.org/10.1007/978-981-99-7962-2_30
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>

- Panfili, L., Duman, S., Nave, A., Ridgeway, K. P., Eversole, N., & Sarikaya, R. (2021). Human-AI interactions through a Gricean lens. *Proceedings of the Linguistic Society of America*, 6(1), 288–302. <https://doi.org/10.3765/plsa.v6i1.4971>
- Qian, Y. (2025). Prompt engineering in education: A systematic review of approaches and educational applications. *Journal of Educational Computing Research*, 63(7–8), 1782–1818. <https://doi.org/10.1177/07356331251365189>
- Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., & Sutskever, I. (2019). Language models are unsupervised multitask learners. *OpenAI Blog*.
- Tour, E., & Zadorozhnyy, A. (2025). Conceptualizing and operationalizing prompt literacy for English language learners. *Journal of Adolescent & Adult Literacy*, 69(3), e70020. <https://doi.org/10.1002/jaal.70020>
- Wei, J., Wang, X., Schuurmans, D., Bosma, M., Ichter, B., Xia, F., Chi, E., Le, Q., & Zhou, D. (2022). Chain-of-thought prompting elicits reasoning in large language models. *arXiv*. <https://doi.org/10.48550/arXiv.2201.11903>
- Zhou, Y., Muresanu, A. I., Han, Z., Paster, K., Pitis, S., Chan, H., & Ba, J. (2023). Large language models are human-level prompt engineers. In *Proceedings of the Eleventh International Conference on Learning Representations (ICLR 2023)*.