

WHAT MAKES AI PROMPTS POPULAR? EXAMINING LISTING PRESENTATION CUES IN AI PROMPT MARKETPLACES

Yuzhang Han, College of Business Administration

California State University San Marcos, San Marcos, United States

Li Sun, College of Science, Technology, Engineering, and Mathematics

California State University San Marcos, San Marcos, United States

Zhuyun Zhang, North Carolina State University, Raleigh, United States

ABSTRACT

AI prompt marketplaces are emerging digital environments where prompts are created, displayed, evaluated, and sometimes sold as market-facing digital products. Despite the growing practical interest in these marketplaces, marketing research has paid little attention to this environment or, particularly, to what makes prompts popular within a marketplace. Situated in prompt marketplaces, this study examines prompt popularity, an underexplored marketplace response construct, and focuses on how it relates to the design of a prompt's listing page. Drawing on 3,200 listing pages from four representative AI prompt marketplaces, we examine five types of listing page presentation cues: simplicity, symbolic markup, directiveness, specificity, and creator credibility. The analysis uses fixed-effects ordinary least squares (OLS) regression, checked for robustness with within-platform standardized outcome. Our results suggest that prompt popularity is higher when prompt titles are designed to be shorter, less abstract, and less sentimental, and when prompt introductions appear more readable, use symbolic markup, and avoid command-style wording. In practice, this study identifies listing page presentation design as an actionable pathway for platform users and operators to improve prompt value communication, marketplace discovery, and user experience.

Keywords: AI Prompt Marketplaces; Prompt Popularity; Listing Presentation Cues; Digital Marketing; E-Commerce; Generative AI

INTRODUCTION

Generative AI systems have been adopted quickly (Bick et al., 2026). Users interact with these systems through prompts, the natural-language instructions AI follows to produce desired output. Because a model's response can change greatly depending on how the prompt is written, effective prompt design has become an important skill for AI users (P. Liu et al., 2023). However, not everyone has the time or expertise to craft prompts well. As a result, a new intermediary has emerged: the AI prompt marketplace, an online platform where prompts are created, displayed, searched, evaluated, and sometimes sold and bought as reusable digital products.

These prompt marketplaces are driven by the interaction of three actors. Platform users look for prompts that can help them complete tasks such as writing, coding, and image generation; prompt creators develop and sometimes sell prompts; and marketplace platforms themselves organize the prompt inventory and connect the other two actors. The global prompt marketplace had grown to roughly USD 1.3 billion in 2024 (AI Prompt Marketplace Market, 2025). The sector is expected to expand at 29.5% annually, reaching approximately USD 7 billion by 2030 (AI Prompt Marketplace Market Report 2026, 2026).

On prompt marketplaces, prompt popularity has become a key measure of prompt success. It refers to the visible positive responses from users a prompt accumulates, displayed on the prompt's listing page as popularity score, upvotes, favorites, or likes, depending on the platform. For users, prompt popularity acts as social proof that signals which prompts are worth trying, as the popularity score does for other digital products (Salganik et al., 2006). For creators, it may raise the visibility of their prompts, helping their work get discovered by more users. For platforms, it may influence how prompts are ranked and recommended to users. Prompt popularity is difficult to study, because a prompt's full value is rarely evident from the listing page alone; it becomes clearer only after users apply the prompt to their own tasks.

Despite the practical importance of prompt marketplaces and prompt popularity within these platforms, research on this emerging setting has three gaps. First, AI prompt marketplaces are themselves underexplored. Marketing and information systems research has examined adjacent settings such as app stores (Ghose & Han, 2014), e-commerce and online retail websites (Jawale et al., 2026; Jiang & Benbasat, 2007), and social platforms (Berger & Milkman, 2012; Sharma et al., 2026) yet it has paid only limited attention to prompt marketplaces as a distinct environment. Second, even within the emerging research on prompt marketplaces, prompt popularity has rarely been the focus. Existing work centers on prompt pricing (M. Li et al., 2024) and sales (Cao et al., 2024), platform risks (Hou et al., 2025) and prompt security (Shen et al., 2024; Trinh et al., 2024, 2026; Wu et al., 2025; Zou, 2025), and platform design (X. Li et al., 2024). By contrast, prompt popularity has received very limited research attention.

Together, these limitations lead to a third gap: prompt popularity has not been studied in relation to listing presentation cues. We use this term to refer to observable cues displayed on a prompt's listing page, such as its wording, structure, formatting, and creator-side information, that are independent of the type, topic, or task domain of the underlying prompt. Studying these cues is especially meaningful because they influence how users evaluate a prompt before they can test it, when the prompt's true value is still uncertain. Because these cues do not depend on the underlying prompt, they are comparable across prompts. Findings about these cues may therefore generalize across prompt categories, paid and free platforms, and potentially other digital marketplaces. They can guide creators and platform designers regardless of what a given prompt is about.

Beyond prompt marketplaces, prior research shows that presentation cues are linked to user response and to the outcomes of digital offerings. For example, in e-commerce, the cues on a seller's website are associated with product quality (Mavlanova et al., 2012). In crowdfunding, textual and linguistic cues in a campaign's listing are associated with its funding success

(Kaminski & Hopp, 2020). However, presentation cues have not yet been examined in relation to prompt popularity.

This study addresses these gaps by treating prompt listing pages as market-facing interfaces for digital products and asking: *How are listing presentation cues associated with prompt popularity?* To answer this question, we collected 3,200 prompt pages from four representative prompt marketplaces. Using these data, we examine how listing page cues are associated with prompt popularity displayed on platforms. These cues include listing page simplicity, specificity, directiveness, use of symbolic markup, and creator information.

The study makes several contributions. It investigates prompt marketplaces, an underexplored environment for marketing research; it also explores and operationalizes prompt popularity, an underexamined marketplace response construct. Furthermore, it examines how observable, content-independent features of prompt listing pages are associated with prompt popularity. Findings of the study offer guidance for prompt creators seeking to communicate prompt value more effectively and for marketplace operators seeking to design platform interfaces, ranking systems, and listing standards that improve discovery, trust, and user decision making.

The remainder of the paper is organized as follows. Section Conceptual Background and Hypotheses develops the conceptual framework and hypotheses based on a literature review. Section Methodology describes the data and analytic models. Results presents the analytic results and findings. Discussion and Implications sections interpret the results and discuss the study's relevance. Finally, Section Conclusion, Limitations, and Future Work concludes the paper and outlines future research directions.

CONCEPTUAL BACKGROUND AND HYPOTHESES

Prompt Marketplaces as Emerging Digital Marketplace Environment

A prompt marketplace is a multi-sided platform that connects prompt creators and users and operates through cross-side network effects (Rochet & Tirole, 2003). It resembles app stores, online retail platforms, and content platforms because it lists, organizes, and ranks items supplied by external creators (Ghose & Han, 2014). However, the traded good differs from more established marketplace objects. A prompt is an instruction rather than a finished product, and its quality is difficult to judge before use. A prompt usually has to be tested in a generative AI model before users can see what it actually produces (P. Liu et al., 2023). Even then, the result may vary from one run to another because most AI models are probabilistic (Atil et al., 2024). Therefore, users cannot fully know a prompt's usefulness in advance. Before trying it, they often have to evaluate the prompt through the visible presentation cues on the listing page.

Research on prompt marketplaces is still at an early stage and spread across different topics. Some studies examine the commercial mechanism of these markets, such as how prompts can be priced (M. Li et al., 2024). Some other studies inspect marketplace security, focusing on prompt stealing and platform vulnerabilities (Hou et al., 2025; Shen et al., 2024; Wu et al., 2025; Zou, 2025). Another line of studies looks at the relationship between prompts and the content generated across model types, including community-built chatbots (X. Li et al., 2024) and

artwork generators (Trinh et al., 2024, 2026). Most similar to our research, (Cao et al., 2024) links linguistic and demonstration cues on prompt listing pages to prompt sales. The present study extends this listing cue perspective to prompt popularity.

This study draws on research from related platform settings. In e-commerce and online retail, prior studies show that product presentation can influence perceived quality and demand (Jiang & Benbasat, 2007; Mavlanova et al., 2012). Crowdfunding research also shows that textual features of a listing, such as language style and readability, are related to funding success (Escudero et al., 2026; Kaminski & Hopp, 2020). Work on social and content platforms further shows that content characteristics can affect popularity and virality (Berger & Milkman, 2012; Zahrah et al., 2026). Building on these areas, this study examines how listing page cues matter in the newer context of prompt marketplaces.

Prompt Popularity as Marketplace Response Construct

In this study, prompt popularity refers to the number of endorsements a prompt receives from users. It is displayed on the listing page as the count of favorites, likes, upvotes, or the popularity score. As a marketplace response construct, prompt popularity acts as social proof of a prompt's quality (Salganik et al., 2006).

Prompt popularity differs from other marketplace indicators in several ways. Prompt sales and price are tied to monetization and financial decisions (Cao et al., 2024; M. Li et al., 2024). Differently, popularity is non-monetary and appears on both free and paid platforms. Prompt views, downloads, or usage counts only show that a prompt was viewed or tried (Ghose & Han, 2014). Prompt popularity also reflects a user's deliberate choice to endorse the prompt. Furthermore, some studies analyze prompts based on the model output they produce (Trinh et al., 2024, 2026). By contrast, prompt popularity captures how users themselves judge the prompt. Overall, Prompt popularity is both a reflection of prior user responses and a factor that impacts future user responses.

Using popularity as a dependent variable is consistent with prior research that studies similar forms of user response in digital markets. Studies have examined post popularity and virality on cultural and social platforms (Berger & Milkman, 2012; Salganik et al., 2006; Yang et al., 2026), product demand and ranking in app stores (Ghose & Han, 2014), as well as funding success in crowdfunding (Kaminski & Hopp, 2020). Across these settings, researchers often link visible listing or content cues to user response. This study applies the same general approach to prompt popularity in prompt marketplaces.

Listing Presentation Cues as Marketplace Signals

Signaling theory suggests that when information is unevenly held, a better-informed party can send visible signals, and a less-informed party uses them to infer qualities it cannot observe directly (Connelly et al., 2011; Spence, 1978). This idea fits prompt marketplaces because users cannot fully know a prompt's full value before trying it. Before use, they often judge the prompt using information shown on its listing page. In this study, we treat the listing page as a set of visible signals. We define listing presentation cues as visible page features that are independent of the underlying prompt, such as page wording, formatting, and creator information, that can be

observed without accessing and testing the prompt itself. These cues can be compared across prompts even when the prompts cover different tasks or topics.

Prior research in other digital settings shows that visible cues can help explain market outcomes. In online retail, product presentation and visible signals have been linked to buyers' perceived quality and product demand (Jiang & Benbasat, 2007; Mavlanova et al., 2012; X. Wang & Xu, 2026). A similar pattern appears in crowdfunding, where the textual and linguistic features of a campaign are associated with funding success (Kaminski & Hopp, 2020). On content platforms, research shows that post characteristics help explain its virality (Berger & Milkman, 2012). In prompt marketplaces, linguistic and demonstration signals on prompt listing pages are related to prompt sales (Cao et al., 2024). Our study extends this line of research by relating listing presentation cues to a new response, prompt popularity.

Hypothesis Development and Theoretical Lenses

Under the signaling framework, we develop five hypotheses drawing on different theoretical lenses, each linking a different type of listing cue to prompt popularity.

Processing fluency suggests that people respond more positively to information that is easy to read and understand (Alter & Oppenheimer, 2009). In prompt marketplaces, this means that simpler listing pages may help users understand a prompt more quickly and with less effort, which may increase prompt popularity. For example, a short, plain prompt title can make the prompt easier to scan, while a readable and syntactically simple prompt introduction can make its purpose easier to understand. For this reason, listing pages with shorter titles, shorter title words, more readable introductions, and syntactically less complex introductions should be associated with higher prompt popularity.

H1 Simplicity: Prompt listing pages with simpler, easier-to-read wording, reflected in shorter titles, shorter title word length, higher introduction readability, and lower introduction syntactic complexity, are positively associated with prompt popularity.

Symbolic markup can also raise processing fluency. By giving its text components a clear structure, a marked-up listing page is easier to process and should draw a more favorable response. For example, listing pages that use special tokens in the prompt title or non-textual special tokens in the prompt introduction should be easier to process and therefore more likely to gain popularity.

H2 Markup: Prompt listing pages that use more symbolic markup, reflected in a higher proportion of special tokens in the title and a higher proportion of symbols in the introduction, are positively associated with prompt popularity.

Research on persuasion knowledge and psychological reactance suggests that overtly persuasive or command-oriented language can make users more resistant to a message. In prompt marketplaces, listing pages that rely heavily on command-style language may be perceived less as product descriptions and more as direct instructions. This can make users more cautious and less willing to respond positively. For this reason, for example, listing pages with more imperative language in the introduction or more verbs in the title should be associated with lower prompt popularity.

H3 Directiveness: Prompt listing pages with more command- and action-oriented language, reflected in a higher proportion of imperative verbs in the introduction and a higher proportion of verbs in the title, are negatively associated with prompt popularity.

Research on message concreteness suggests that readers respond more favorably when a message is specific; compared with abstract wording, concrete wording attracts more attention, appears more important, and improves evaluations of the message source (Miller et al., 2007). Consumer research further shows that consumers are less skeptical of objective claims than subjective claims, partly because objective claims are more precise and easier to verify (Ford et al., 1990). On this basis, concrete, specific, and objective listing page should be associated with higher prompt popularity. For example, a less abstract title or a more concrete introduction may help users judge what a prompt can offer, while an overly sentimental title may sound exaggerated and weaken users' willingness to endorse it.

H4 Specificity: Prompt listing pages with more specific and concrete wording, reflected in fewer abstract tokens in the title, less positive title sentiment, and higher introduction concreteness, are positively associated with prompt popularity.

Finally, source credibility theory holds that messages from a more expert, trustworthy source are more readily accepted (Hovland & Weiss, 1951; Ismagilova et al., 2020). On this basis, listings whose creator information conveys greater credibility should be more popular. For example, a creator who has produced many prompts has a track record that signals expertise, while a prompt with fewer creators has a clearer, more accountable source; both should make the listing more popular.

H5 Creator: Prompt listings with creator-side signals of higher source credibility, reflected in a higher creator prompt count and a lower creator count, are positively associated with prompt popularity.

METHODOLOGY

Research Design

We examine the relationship between listing presentation cues and prompt popularity using data from four AI prompt marketplaces. The sample includes two paid platforms (prompts must be purchased before they can be accessed in full) and two free platforms (prompts can be accessed in full without payment). To select the platforms, we reviewed dozens of possible candidates and chose platforms that met four criteria well established in platform research and online data collection methodology. The platforms needed to have enough scale or prompt resources (Haim et al., 2018; McIntyre & Srinivasan, 2017), well-structured listing pages and platform interface (Mancosu & Vegetti, 2020), sufficient visibility in industry or scholarly discussion (Rietveld & Schilling, 2021), and meaningful variation across platform types (Seawright & Gerring, 2008).

Each selected marketplace is a prominent example within its segment for scale, listing interface clarity, and visibility. They are: PromptBase (paid) is the most widely cited marketplace and an early mover in monetized prompt exchange, with a highly structured listing interface (Cao et al., 2024; M. Li et al., 2024; PromptBase, 2026; Trinh et al., 2024, 2026; Zhang et al., 2025); LaPrompt (paid) is another leading marketplace, with verified prompt resource and a rich, category-organized platform interface (LaPrompt, 2026); FlowGPT (free) is the most-covered free platform in the press, with the richest platform interface and a very large user base (X. Li et al., 2024); and prompts.chat (free) is the largest open-source, community-driven

platform, with noticeable academic visibility and a prominent GitHub origin.(prompts.chat, 2026; Zhang et al., 2025).

Data Collection and Predictor Selection

Prompt marketplaces increasingly restrict or even ban automated data collection for data security reasons. Therefore, we had to collect the data manually. While manual collection helped ensure compliance with platform policies, it also limited the number of platforms and the sample size that could be collected. During data collection, two collectors gathered the 800 highest-ranked listing pages from each platform, ordered by that platform's popularity indicator (described in Section Measures). Collection was completed in February 2026 and took about three to five days for each platform.

Using top ranked or high ranked listings as the item pool is a common practice in online platform research. This approach is especially appropriate when the research focuses on how publicly visible listing information relates to marketplace outcomes, such as search visibility and popularity (e.g., Jozani et al., 2025; Jürgensmeier & Skiera, 2025). This sampling strategy fits our study because top-ranked prompt pages are more likely to be carefully designed and maintained, making their listing presentation cues more meaningful for analysis. At the same time, it reduces the inclusion of inactive, incomplete, experimental, or low-effort pages, which are especially common on free platforms and often do not contain well-developed presentation features.

We selected predictors based on their interpretability, conceptual coverage, data availability, and low collinearity. From the collected listing pages, we derived an initial pool of more than 200 candidate features. We first screened these for data availability ($\geq 95\%$ non-missing) and variability (≥ 2 unique values), reducing the pool to 47. The remaining variables entered a greedy forward selection procedure (Heinze et al., 2018), with candidates considered in order of their theoretical importance established in prior marketplace research (e.g., Ghose & Ipeirotsis, 2010). To limit collinearity, a candidate predictor was retained only if its absolute pairwise correlation with all selected predictors stayed below 0.65, yielding a final set of 13 predictors.

Measures

The dependent variable, prompt popularity, is implemented by the positive endorsement count shown on each listing page. It is operationalized as the favorite count on PromptBase, like count on LaPrompt, popularity score on FlowGPT, and upvote count on prompts.chat. The prompts.chat platform does not provide a “downvote” option. Therefore, its upvote count captures positive endorsements only and is comparable to the other platform indicators.

The independent variables include 13 predictors that represent five types of listing presentation cues: simplicity (H1), markup (H2), directiveness (H3), specificity (H4), and creator information (H5). Two exploratory interaction terms are also included to examine how key cue types interact. Their definitions are presented in Table A1 in Appendix A.

Data Description

The analytical sample includes 3,200 listing pages, with 800 observations from each of the four platforms. Table 1 reports the sample structure and distributions of prompt popularity at the platform level. We can see that prompt popularity varies substantially across platforms. Overall, free platforms show higher popularity than paid platforms. FlowGPT stands out from the other three platforms. This pattern is shown in Figure 1, where FlowGPT's popularity distribution lies clearly above the others.

These large platform differences are expected and are part of the study design. Prompt marketplaces vary widely in monetization model, platform scale, interface structure, user base, and the type of prompt resources they host. Therefore, we selected a heterogeneous set of platforms to surface this diversity, allowing the findings to generalize beyond a single platform type. The resulting popularity gaps likely reflect these platform-level differences. We treat such popularity gaps as a feature to be incorporated rather than noise to be removed. Section Model Specification accounts for the gaps through log-transformed popularity, platform fixed effects, and within-platform robustness checks.

Scope	Sample Size	Mean Popularity	Median Popularity	SD Popularity	Max Popularity	Mean Popularity
FlowGPT (free)	800	39,507,932.250	14,350,000.000	73,535,785.708	686,800,000.000	15.110
prompts.chat (free)	800	0.463	0.000	1.178	14.000	0.232
PromptBase (paid)	800	46.586	35.000	42.669	726.000	3.618
LaPrompt (paid)	800	0.785	0.000	1.438	16.000	0.384
Free Platforms	1600	19,753,966.356	3,107.000	55,610,504.346	686,800,000.000	7.671
Paid Platforms	1600	23.686	6.000	37.889	726.000	2.001

Note. Values are reported for the balanced analytical sample. Popularity is the platform-displayed count of positive user endorsements on the prompt page; the transformed outcome is $\ln(1 + \text{popularity})$.

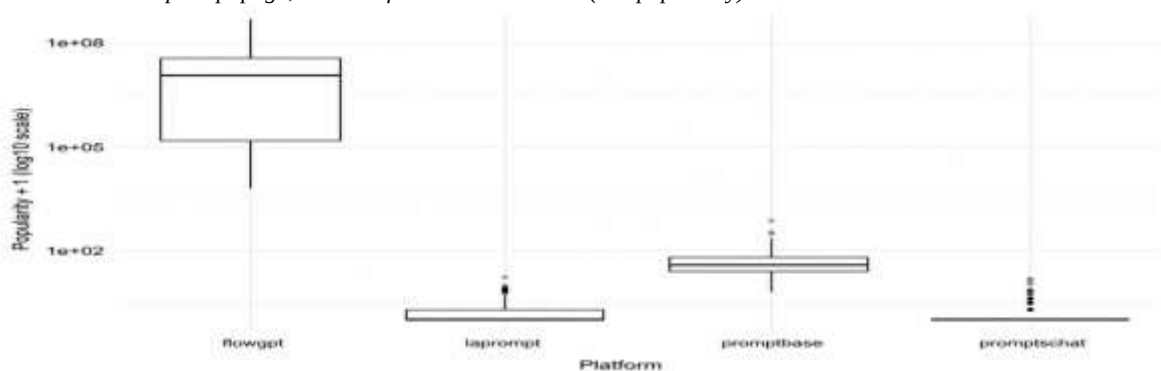


FIGURE 1 PLATFORM-LEVEL DISTRIBUTION OF PROMPT POPULARITY

Note. The figure plots popularity + 1 on a log10 scale for the analytical sample, with 800 prompt pages per platform.

Model Specification

As shown in Equation 1, the main model is estimated using Ordinary Least Squares (OLS) with platform fixed effects (Mundlak, 1978). Log transformation is applied to the outcome variable, prompt popularity.

$$y_{ij} = \ln(1 + [\text{popularity}]_{ij}) = \alpha_j + X_{ij} \beta + Z_{ij} \theta + \varepsilon_{ij} \quad (1)$$

For prompt i on platform j , X_{ij} holds the main-effect predictors, Z_{ij} the interaction terms, α_j the platform fixed effects, and ε_{ij} the residual. The log transform of prompt popularity compresses the right-skewed outcome and reduces the influence of differences in popularity across platforms (Curran-Everett, 2018). α_j further absorb systematic differences in baseline popularity across platforms, allowing the results to focus on within-platform associations between listing page cues and prompt popularity.

An additional model serves as robustness check, shown in Equation 2. The model re-estimates the main model on the outcome standardized within platform (L. Wang et al., 2019). This approach tests whether the findings are robust to cross-platform differences in both baseline popularity and the dispersion of popularity values within each platform.

$$y_{ij}^z = (\ln(1 + [\text{popularity}]_{ij}) - \bar{y}_j) / s_j = \alpha_j + X_{ij} \beta + Z_{ij} \theta + \varepsilon_{ij} \quad (2)$$

Here, $\bar{y}_j$ and s_j are the mean and standard deviation of the log-transformed popularity. y_{ij}^z expresses how far prompt i 's log-popularity falls above or below its platform's mean, in units of that platform's standard deviation.

Comparing these two models, Equation 1 serves as the main model because its coefficients are on the log-popularity scale, which matches how the hypotheses are stated.

Furthermore, multiple listing pages may be created by the same creator. To account for the possible dependence among pages from the same creator, we use the creator IDs shown on listing pages to define creator clusters. These clusters are then used to estimate robust standard errors used in all models (MacKinnon et al., 2023).

RESULTS

Model Diagnostics

Table 2 reports the diagnostic results for the fixed-effects OLS model. Overall, the model shows high explanatory power, with an R^2 of 0.937 (Gao, 2024). Multicollinearity appears limited, supported by the largest variance inflation factor (VIF) adjusted for fixed effects of 1.630 (Tatarynowicz & Keil, 2026). The Breusch-Pagan test is significant ($\chi^2 = 1358.009$, $p < 0.001$), indicating heteroskedasticity (Breusch & Pagan, 1979). The raw popularity count is overdispersed, with a variance-to-mean ratio of 1.665×10^8 (del Castillo & Pérez-Casany, 2005). Finally, the intra-cluster correlation coefficient (ICC) is high, at 0.992 across 1,410 creator clusters. This indicates strong dependence among listing pages from the same creator and

motivates the use of robust standard errors clustered by creators (Killip et al., 2004; MacKinnon et al., 2023).

TABLE 2 MODEL DIAGNOSTICS	
Metric	Value
Complete-case observations (main model)	3,191.000
Main-effect count	13.000
Interaction count	2.000
Modeled term count	15.000
Creator clusters	1,410.000
ICC for $\ln(1 + \text{popularity})$	0.992
Main FE OLS R^2	0.937
AIC, main FE OLS	12035.699
BIC, main FE OLS	12150.993
Overdispersion ratio (variance/mean)	1.665×10^8
Breusch-Pagan test statistic	1358.009
Breusch-Pagan p value	< 0.001
Largest FE-adjusted VIF	1.630 (intro readability)

Note. Diagnostics are reported for the retained complete-case estimation sample. The Breusch-Pagan test was computed from the equivalent dummy-variable OLS formulation of the fixed-effects model. FE-adjusted VIFs were computed after absorbing platform fixed effects. The overdispersion ratio is the variance-to-mean ratio of the raw popularity count.

Figure 2 reports the pairwise correlations among model terms after platform fixed effects are absorbed. Most correlations are weak to modest. Stronger associations appear mainly among related prompt introduction measures, such as dependency distance, readability, and symbol proportion. This pattern is reasonable, since these variables all describe related aspects of how an introduction reads. Overall, the figure does not suggest serious redundancy among the predictors, which is consistent with the low VIF reported in Table 2.

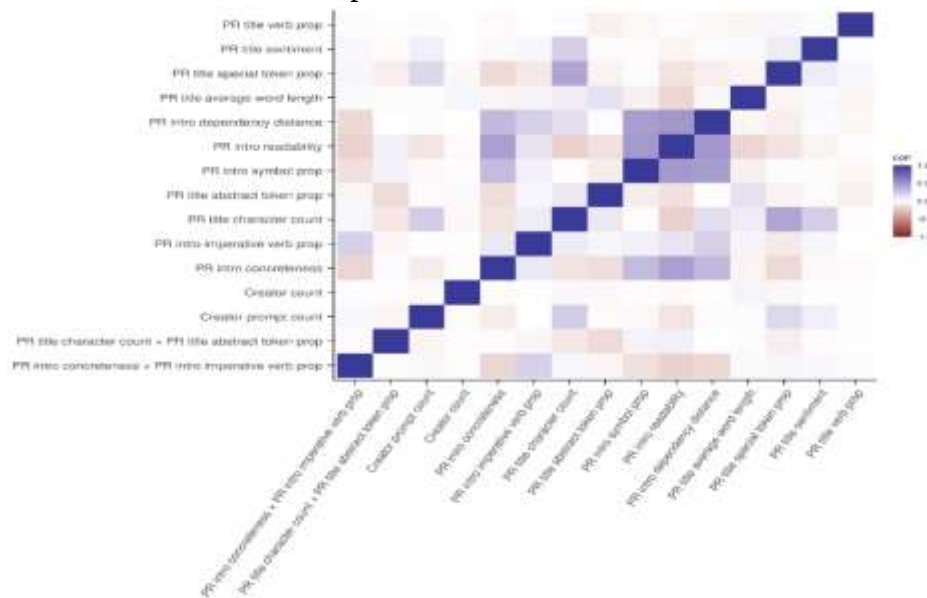


FIGURE 2 PAIRWISE CORRELATIONS AMONG MODEL TERMS

Note. Heatmap of Pearson correlations among all model terms after absorbing platform fixed effects. PR denotes Prompt.

Hypothesis Testing

Table 3 reports the main fixed-effects OLS results for Equation 1. H1 predicted that listing pages designed simpler and easier to read would be positively associated with prompt popularity. The results mostly support H1. For prompt titles, both representative predictors, character count and average word length, are negative and significant, indicating that shorter titles and simpler title wording are associated with higher prompt popularity. For prompt introductions, readability is positive and significant, suggesting that clearer introductions also receive stronger marketplace response. Introduction dependency distance, an indicator of syntactic complexity (Futrell et al., 2015; H. Liu, 2008), is negative as expected but not significant. These findings suggest that listings requiring less reading effort tend to be processed more fluently and, in turn, may be evaluated more favorably. Thus, H1 receives substantial support in the main model, supported by title simplicity and introduction readability.

TABLE 3 MAIN MODEL RESULTS					
Hypothesis	Predictor	Coefficient	SE	p	Std. β
H1 Simplicity	title character count	-0.0149***	0.0041	= .0003	-0.0384
	title avg word length	-0.0997***	0.0266	= .0002	-0.0240
	intro dependency distance	-0.0446	0.0568	= .4323	-0.0069
	intro readability	0.0136***	0.0022	< .0001	0.0648
H2 Markup	title special token prop	3.6807**	1.2019	= .0022	0.0416
	intro symbol prop	3.7228*	1.6837	= .0272	0.0156
H3 Directiveness	title verb prop	-1.6233***	0.4455	= .0003	-0.0277
	intro imperative verb prop	-1.1130***	0.2396	< .0001	-0.0377
H4 Specificity	title abstract token prop	-0.9986***	0.2563	= .0001	-0.0202
	title sentiment	-0.4988**	0.1662	= .0027	-0.0173
	intro concreteness	-0.5885	0.3014	= .0511	-0.0168
H5 Creator	creator count	-0.4330	0.3850	= .2609	-0.0054
	creator prompt count	0.0003	0.0002	= .0943	0.0070
Exploratory	title character count $\times$ title abstract token prop	0.0380*	0.0161	= .0181	0.0108
	intro concreteness $\times$ intro imperative verb prop	0.1931	1.1524	= .8670	0.0013

Note. Coefficients are from the main fixed-effects OLS model with $\ln(1 + \text{popularity})$ as the dependent variable and creator-clustered robust standard errors. Standardized betas are from the corresponding standardized fixed-effects model. Asterisks indicate statistical significance ($p < .05$, ** $p < .01$, *** $p < .001$).*

H2 predicted that the use of symbolic markup in listing pages would be positively associated with prompt popularity. The results support this hypothesis. Both representative predictors, title special token proportion and introduction symbol proportion, are positive and significant. These findings suggest that using markup may help structure listing pages and improve page scannability, making key information easier to identify. As a result, listing pages may receive more favorable user responses.

H3 predicted that more command- and action-oriented language in listing pages would be negatively associated with prompt popularity. The results support H3. Both title verb proportion and introduction imperative verb proportion are negative and significant. This pattern suggests that, although prompts themselves are instructions, listing pages that sound too directive may be less appealing; users may appear to respond more favorably to listings with less command-like language.

H4 predicted that more specific and concrete wording in listing pages would be positively associated with prompt popularity. The results offer partial support for this hypothesis. Title abstract token proportion is negative and significant, suggesting that less abstract prompt titles may make the prompt's value appear more concrete and credible, thus rendering the prompt more popular. Title sentiment is also negative and significant, suggesting that overly positive titles may be perceived as exaggerated or less credible, making the prompt less popular. Introduction concreteness is negative, opposite to the predicted direction, and not significant. Therefore, H4 is supported for prompt title wording, through lower abstractness and more restrained sentiment, but not for the prompt introduction.

H5 predicted that creator-side credibility signals from listing pages would be positively associated with prompt popularity. The results do not support H5. Creator prompt count is positive as predicted, and creator count is negative as predicted. However, neither link reaches statistical significance. These results suggest that creator-side signals may have the expected directional association with popularity, but they do not provide statistically reliable explanatory power once platform fixed effects and listing page presentation cues are included. Prompt popularity appears to be driven more by how the listing page is designed and presented than by creator information.

The exploratory interaction between title character count and title abstract token proportion adds further nuance to H1 and H4. Originally, title character count hosts the negative effect of longer titles in H1, while title abstract token proportion conveys the negative effect of abstract title wording for H4. The positive interaction indicates that the negative effect of abstract title wording may become weaker when titles are longer. One possible interpretation is that additional title length gives users more context, making abstract wording easier to understand. In this way, the interaction adds to H1 and H4, suggesting that title simplicity and title specificity may operate jointly in influencing prompt popularity.

Robustness Check

To examine whether the findings remain stable after accounting for large cross-platform differences in prompt popularity, we re-estimate the model using the within-platform standardized outcome in Equation 2 (results shown in Table B1 of Appendix B). The core findings remain largely consistent: introduction readability (H1), both markup cues (H2), both directiveness cues (H3), and title abstract token proportion (H4) remain significant and in the same direction as in the main model. These results indicate that the evidence for introduction simplicity (H1), symbolic markup (H2), directiveness (H3), and title specificity (H4) is robust to cross-platform heterogeneity.

Several other indicators show weaker robustness. Title average word length (H1) and title sentiment (H4) remain in the predicted negative direction but no longer reach statistical significance. Title character count (H1) becomes insignificant and near zero. Introduction concreteness (H4) becomes significantly negative, opposite to the predicted direction. Together, these results suggest that the title-based component of H1 and the sentiment and concreteness components of H4 should be interpreted with caution after platform differences in popularity levels and variability are accounted for. Finally, creator cues (H5) remain non-significant, consistent with the main model.

The significant exploratory interaction in the main model, title character count $\times$ title abstract token proportion, does not remain significant and the predicted direction in the robustness model. This suggests that the cue combination, where longer titles may soften the negative effect of abstract wording, is not stable once popularity is standardized within platforms.

DISCUSSION

The results provide a mixed but informative picture of prompt popularity and listing page presentation. Prompts are more popular when the wording of their listing pages is simpler (H1), as reflected in using shorter, plainer prompt titles and more readable prompt introductions, which may make the pages easier to process (title character count: Std. $\beta = -0.0384$, $p = .0003$; title average word length: Std. $\beta = -0.0240$, $p = .0002$; introduction readability: Std. $\beta = 0.0648$, $p < .0001$). Higher prompt popularity is also positively associated with the use of symbolic markup (H2), as reflected in the use of special tokens in titles and symbols in the introductions, which could make key information easier to identify (title special token proportion: Std. $\beta = 0.0416$, $p = .0022$; introduction symbol proportion: Std. $\beta = 0.0156$, $p = .0272$). Specific title wording (H4) shows the same positive effect, as reflected in titles that are less abstract and less emotionally positive; such titles may give users a clearer and more concrete impression of the prompt (title abstract token proportion: Std. $\beta = -0.0202$, $p = .0001$; title sentiment: Std. $\beta = -0.0173$, $p = .0027$).

By contrast, directive wording (H3) has a negative effect, as reflected in command-oriented titles and introductions, which could make the listing appear more forceful and less appealing to users (title verb proportion: Std. $\beta = -0.0277$, $p = .0003$; introduction imperative verb proportion: Std. $\beta = -0.0377$, $p < .0001$). Finally, the creator cues (H5) point in the expected direction, but their effects are not statistically significant (creator count: Std. $\beta = -0.0054$, $p = .2609$; creator prompt count: Std. $\beta = 0.0070$, $p = .0943$).

Several observations follow from these results. First, prompt popularity depends not only on the prompt itself, an underlying technical artifact, but also on how the prompt is presented on its listing page. The results suggest that presentation features matter, including how readable the listing page is (H1), how it is structured with markup (H2), how directive its wording is (H3), and how specific its title is (H4). These cues are part of the information users can process when reading the listing page. They are different from the underlying qualities of the prompt, which users cannot fully observe before use. In this sense, prompt marketplaces resemble other digital product environments, where visible presentation helps users make judgments under uncertainty.

Second, the title functions as an especially important entry point for user evaluation. Various title properties are associated with prompt popularity, including shorter and simpler title wording (H1), as well as lower title abstractness and less strongly positive title sentiment (H4). Title length and title abstractness also interact (title character count $\times$ title abstract token prop: Std. $\beta = 0.0108$, $p = .0181$), and this combined effect is positively associated with prompt popularity. Overall, title design appears most effective when it gives users a concise, clear, and specific first impression of the prompt.

Third, user evaluation in the prompt marketplace setting appears to be more message centered than source centered. Text-based cues in the title and introduction consistently predict prompt popularity (H1 through H4), while creator credibility cues do not (H5). Users appear to judge a prompt from how it is presented rather than from who supplied it, unlike other digital marketplaces such as app stores, where developer reputation drives demand (Ghose & Han, 2014). One reason may be that a prompt is itself a short piece of text, which can be conveyed clearly through its title and introduction. This makes users less reliant on the creator as a signal of quality.

IMPLICATION

For marketing research, this study explores AI prompt marketplaces, a meaningful, underexplored context from the perspectives of digital marketing and electronic commerce. We show that prompts are not only technical instructions; on prompt marketplaces they become market-facing digital products that are evaluated under uncertainty, as their quality is not directly observable before use and must be inferred from the listing page. Furthermore, by focusing on prompt popularity, an underexplored construct, the study extends prompt marketplace research beyond sales, price, and security to a visible endorsement outcome present on almost all types of platforms. By doing this, the study connects prompt marketplace research to broader marketing questions about product presentation, social proof, and platform-mediated consumer evaluation.

For prompt creators, who in effect market their own products, this study demonstrates that a prompt's popularity depends not only on how useful it is, but also on how clearly and credibly its listing page conveys that usefulness. Therefore, the listing page serves as the prompt's primary interface to users, informing their judgment before they use it. Drawing on the study's results, creators can strengthen the page with concise titles, readable introductions, structured markup, and specific wording that let users quickly grasp what a prompt offers, while avoiding the command-oriented or exaggerated phrasing that tends to weaken user response. This matters most for small creators, who may lack established brand recognition and therefore rely on the listing page itself to convey value.

For marketplace operators, this study introduces listing presentation as a manageable factor for strengthening marketplace quality and competitiveness. Operators can raise platform-wide listing page quality by giving page templates and formatting tools that encourage clarity, structure, and specificity. This enables creators to convey their prompts more clearly. Platforms can also incorporate listing presentation qualities as metrics for prompt ranking and recommendation. Well-organized listing pages would then be easier for users to find and quicker to assess. At the marketplace governance level, operators can embed these presentation standards

into platform policy. This helps reduce noisy or exaggerated listings and facilitates a more reliable marketplace.

CONCLUSION, LIMITATIONS, AND FUTURE WORK

This study examined how listing presentation cues are associated with prompt popularity in AI prompt marketplaces. Analyzing 3,200 prompt listing pages from four platforms, our results show several presentation cues associated with higher prompt popularity, including clearer wording, symbolic markup, lower directiveness, and more specific title design. These findings suggest that prompt popularity is influenced strongly by how prompts are presented on marketplace listing pages. This study also offers practical guidance for prompt creators and marketplace operators. It identifies listing presentation design as a tool for improving prompt value communication and marketplace user experience.

Several limitations exist due the scope of the study. First, the predictor set is limited. Our analysis focuses on presentation cues that are observable on listing pages. In turn, factors unrelated to listing presentation design are not considered, such as page age and user reviews. In addition, to stay representative of diverse marketplaces, the study draws on multiple platforms and therefore excludes cues that are not available on any platforms included. This excludes cues unique to paid platforms, such as prompt price and sales, and cues collectable only on free platforms, such as prompt body content features. Second, the analysis examines prompts at a general level, without further inspecting whether the cues operate differently within specific prompt domains, such as coding, marketing, or creative writing. Third, prompt popularity is measured through a single metric, the visible endorsement count. Other comparable signals, such as prompt views and downloads, could be incorporated and compared given more time and analytic resources.

Future research can extend this work in multiple directions. First, to capture signals beyond the textual presentation cues, research can examine other aspects of listing pages and creator background, such as page age, user reviews, and creator performance. Second, future studies can connect cues to a prompt's actual performance after use, measured by model output quality, task success, and user satisfaction. Third, researchers can examine how platform policies, such as business model and data security rules, may interact with listing page cues and user evaluation in prompt marketplaces.

REFERENCE

- AI Prompt Marketplace Market Report 2026*. (2026). <https://www.thebusinessresearchcompany.com/report/artificial-intelligence-ai-prompt-marketplace-market-report>
- AI Prompt Marketplace Market*. (2025). <https://market.us/report/ai-prompt-marketplace-market/>
- Alter, A. L., & Oppenheimer, D. M. (2009). Uniting the tribes of fluency to form a metacognitive nation. *Personality and social psychology review*, 13(3), 219-235.
- Atil, B., Aykent, S., Chittams, A., Fu, L., Passonneau, R. J., Radcliffe, E., ... & Baldwin, B. (2024). Non-determinism of "deterministic" llm settings. *arXiv preprint arXiv:2408.04667*.
- Berger, J., & Milkman, K. L. (2012). What makes online content viral?. *Journal of marketing research*, 49(2), 192-205.
- Bick, A., Blandin, A., & Deming, D. J. (2026). The rapid adoption of generative AI. *Management Science*.

- Breusch, T. S., & Pagan, A. R. (1979). A simple test for heteroscedasticity and random coefficient variation. *Econometrica: J Econometric Soc* 1287–1294.
- Cao, C., Zhao, L., Li, Y., & Xie, C. (2024, May). Selling in Prompt Marketplace: An Empirical Study on the Joint Effects of Linguistic and Demonstration Signals on Prompt Sales. In *Wuhan International Conference on E-business* (pp. 264-275). Cham: Springer Nature Switzerland.
- Connelly, B. L., Certo, S. T., Ireland, R. D., & Reutzel, C. R. (2011). Signaling theory: A review and assessment. *Journal of management*, 37(1), 39-67.
- Curran-Everett, D. (2018). Explorations in statistics: the log transformation. *Advances in physiology education*, 42(2), 343-347.
- del Castillo, J., & Pérez-Casany, M. (2005). Overdispersed and underdispersed Poisson generalizations. *Journal of Statistical Planning and Inference*, 134(2), 486–500.
- Escudero, S. B., Anglin, A. H., Allison, T. H., & Wolfe, M. T. (2026). Crowdfunding: A theory-centered review and roadmap of the multidisciplinary literature. *Journal of Management*, 52(1), 131-184.
- Ford, G. T., Smith, D. B., & Swasy, J. L. (1990). Consumer skepticism of advertising claims: Testing hypotheses from economics of information. *Journal of Consumer Research*, 16(4), 433–441.
- Friestad, M., & Wright, P. (1994). The persuasion knowledge model: How people cope with persuasion attempts. *Journal of consumer research*, 21(1), 1-31.
- Futrell, R., Mahowald, K., & Gibson, E. (2015). Large-scale evidence of dependency length minimization in 37 languages. *Proceedings of the National Academy of Sciences*, 112(33), 10336-10341.
- Gao, J. (2024). R-Squared (R²)—How much variation is explained?. *Research Methods in Medicine & Health Sciences*, 5(4), 104-109.
- Ghose, A., & Han, S. P. (2014). Estimating demand for mobile applications in the new economy. *Management Science*, 60(6), 1470-1488.
- Ghose, A., & Ipeirotis, P. G. (2010). Estimating the helpfulness and economic impact of product reviews: Mining text and reviewer characteristics. *IEEE transactions on knowledge and data engineering*, 23(10), 1498-1512.
- Haim, M., Kümpel, A. S., & Brosius, H. B. (2018). Popularity cues in online media: A review of conceptualizations, operationalizations, and general effects. *Studies in Communication| Media*, 7(2), 186-207.
- Heinze, G., Wallisch, C., & Dunkler, D. (2018). Variable selection—a review and recommendations for the practicing statistician. *Biometrical journal*, 60(3), 431-449.
- Hou, X., Zhao, Y., & Wang, H. (2025, May). On the (in) security of llm app stores. In *2025 IEEE Symposium on Security and Privacy (SP)* (pp. 317-335).
- Hovland, C. I., & Weiss, W. (1951). The influence of source credibility on communication effectiveness. *Public opinion quarterly*, 15(4), 635-650.
- Ismagilova, E., Slade, E., Rana, N. P., & Dwivedi, Y. K. (2020). The effect of characteristics of source credibility on consumer behaviour: A meta-analysis. *Journal of Retailing and Consumer Services*, 53, 101736.
- Jawale, D. S., Singh, J., & Berad, N. R. (2026). Examining the Mediating Role of E-Trust in the Relationship Between E-Utility and E-Loyalty in Online Shopping. *Academy of Marketing Studies Journal*, 30(1), 1–12.
- Jiang, Z., & Benbasat, I. (2007). The Effects of Presentation Formats and Task Complexity on Online Consumers' Product Understanding1. *MIS quarterly*, 31(3), 475-500.
- Jozani, M., Liu, C. Z., Zhu, H., Liu, L., & Choo, K. K. R. (2025). A Network Analysis of Mobile App Top Chart Appearance and Survival. *Information Systems Frontiers*, 1-23.
- Jürgensmeier, L., & Skiera, B. (2025). Measuring self-preferencing on digital platforms. *Journal of Marketing*, 00222429251360772.
- Kaminski, J. C., & Hopp, C. (2020). Predicting outcomes in crowdfunding campaigns with textual, visual, and linguistic signals. *Small Business Economics*, 55(3), 627-649.
- Killip, S., Mahfoud, Z., & Pearce, K. (2004). What is an intracluster correlation coefficient? Crucial concepts for primary care researchers. *The Annals of Family Medicine*, 2(3), 204-208.
- LaPrompt. (2026). *Top AI Prompt Marketplace*. <https://laprompt.com/>
- Li, M., Ren, H., Xiong, H., Qian, Z., & Zhang, X. (2024). Online Prompt Pricing based on Combinatorial Multi-Armed Bandit and Hierarchical Stackelberg Game. *arXiv preprint arXiv:2405.15154*.

- Li, X., Han, Y., Liu, D., An, P., & Niu, S. (2024, November). Flowgpt: Exploring domains, output modalities, and goals of community-generated ai chatbots. In *Companion Publication of the 2024 Conference on Computer-Supported Cooperative Work and Social Computing* (pp. 355-361).
- Liu, H. (2008). Dependency distance as a metric of language comprehension difficulty. *Journal of Cognitive Science*, 9(2), 159-191.
- 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.
- MacKinnon, J. G., Nielsen, M. Ø., & Webb, M. D. (2023). Cluster-robust inference: A guide to empirical practice. *Journal of Econometrics*, 232(2), 272-299.
- Mancosu, M., & Vegetti, F. (2020). What you can scrape and what is right to scrape: A proposal for a tool to collect public Facebook data. *Social Media+ Society*, 6(3), 2056305120940703.
- Mavlanova, T., Benbunan-Fich, R., & Koufaris, M. (2012). Signaling theory and information asymmetry in online commerce. *Information & management*, 49(5), 240-247.
- McIntyre, D. P., & Srinivasan, A. (2017). Networks, platforms, and strategy: Emerging views and next steps. *Strategic management journal*, 38(1), 141-160.
- Miller, C. H., Lane, L. T., Deatrick, L. M., Young, A. M., & Potts, K. A. (2007). Psychological reactance and promotional health messages: The effects of controlling language, lexical concreteness, and the restoration of freedom. *Human communication research*, 33(2), 219-240.
- Mundlak, Y. (1978). On the pooling of time series and cross section data. *Econometrica: journal of the Econometric Society*, 69-85.
- PromptBase. (2026). Prompt Marketplace. <https://promptbase.com/prompts.chat>.
- Reber, R., Schwarz, N., & Winkielman, P. (2004). Processing fluency and aesthetic pleasure: Is beauty in the perceiver's processing experience?. *Personality and social psychology review*, 8(4), 364-382.
- Rietveld, J., & Schilling, M. A. (2021). Platform competition: a systematic and interdisciplinary review of the literature. *Journal of Management*, 47(6), 1528-1563
- Rochet, J. C., & Tirole, J. (2003). Platform competition in two-sided markets. *Journal of the european economic association*, 1(4), 990-1029.
- Salganik, M. J., Dodds, P. S., & Watts, D. J. (2006). Experimental study of inequality and unpredictability in an artificial cultural market. *science*, 311(5762), 854-856
- Seawright, J., & Gerring, J. (2008). Case selection techniques in case study research: A menu of qualitative and quantitative options. *Political research quarterly*, 61(2), 294-308.
- Sharma, P., Kumar, A., Singh Thakur, S., & Sharda, P. (2026). A Qualitative Comparative Analysis of Paid, Owned, and Earned Media: Key Attributes Shaping Consumer Brand Attitude on Social Media. *Academy of Marketing Studies Journal*, 30(1), 1-16.
- Shen, X., Qu, Y., Backes, M., & Zhang, Y. (2023). Prompt stealing attacks against text-to-image generation models.
- Spence, M. (1978). Job market signaling. In *Uncertainty in economics* (pp. 281-306). Academic Press.
- Tatarynowicz, A., & Keil, T. (2026). Tertius Dolens: Interalter Conflict and Its Negative Impact on Broker Performance. *Organization Science*.
- Trinh, K., Seidenberger, S., Spracklen, J., Wijewickrama, R., Viswanath, B., Jadliwala, M., & Maiti, A. (2026, April). Prompt and Circumstances: Evaluating the Efficacy of Human Prompt Inference in AI-Generated Art. In *International Conference on Computational Intelligence in Music, Sound, Art and Design (Part of EvoStar)* (pp. 145-160). Cham: Springer Nature Switzerland.
- Trinh, K., Spracklen, J., Wijewickrama, R., Viswanath, B., Jadliwala, M., & Maiti, A. (2024). Promptly yours? a human subject study on prompt inference in ai-generated art. *arXiv preprint arXiv:2410.08406*.
- Wang, L., Zhang, Q., Maxwell, S. E., & Bergeman, C. S. (2019). On standardizing within-person effects: Potential problems of global standardization. *Multivariate Behavioral Research*, 54(3), 382-403.
- Wang, X., & Xu, Z. (2026). Viewing inside or outside? The effect of internal detail views of food on consumer trust and purchase intention. *Journal of Retailing and Consumer Services*, 90, 104698.

- Wu, Y., Mu, F., Zhang, Q., Zhao, J., Xu, X., Mei, L., ... & Wang, Y. (2025, July). Vulnerability of text-to-image models to prompt template stealing: A differential evolution approach. In *Findings of the Association for Computational Linguistics: ACL 2025* (pp. 16903-16916).
- Yang, Q., Shen, M., Jiale, H., & Meng, L. M. (2026). Fostering consumer engagement with expressive certainty in entrepreneur-generated content. *European Journal of Marketing*, 60(4), 856-890.
- Zahrah, F. A., Dhenabayu, R., Rahman, M. F. W., Dewi, R. S., & Aminah, Z. H. (2026). Predicting Social Media Post Engagement and Virality Using Graph Neural Network Approaches and Content-Based Features. *Kinetik: Game Technology, Information System, Computer Network, Computing, Electronics, and Control*.
- Zhang, Y., Lin, Y., Khan, A., & Wan, H. (2025). Large Language Model Prompt Datasets: An In-depth Analysis and Insights. *arXiv preprint arXiv:2510.09316*.
- Zou, X. (2025). Reinforcement Learning-Based Prompt Template Stealing for Text-to-Image Models. *arXiv preprint arXiv:2510.00046*.

APPENDIX A. PREDICTOR DETAIL

Table A1 Predictor Definition		
Hypothesis	Predictor	Definition
H1 Simplicity	title character count	Number of characters in prompt title.
	title avg word length	Average word length of prompt title.
	intro readability	Flesch Reading Ease score of prompt introduction. Higher values indicate easier-to-read text.
	intro dependency distance	Mean absolute distance between each token and its syntactic head (the word it grammatically depends on) in prompt introduction. Higher values indicate longer-range syntactic links and greater syntactic complexity or processing difficulty (Futrell et al., 2015; H. Liu, 2008).
H2 Markup	title special token prop	Proportion of tokens containing symbols in the prompt title, including: (1) bracketed or braced tokens that mark structured elements, such as replaceable fields (e.g., {topic} or {product}), user-provided inputs (e.g., [enter your name here]), and standardized input placeholders (e.g., [URL] and [EMAIL]); and (2) special-purpose tokens that contain special characters, such as hashtags (#) and user references (@).
	intro symbol prop	Proportion of non-alphanumeric, non-space characters in prompt introduction.
H3 Directiveness	title verb prop	Proportion of verb in prompt title
	intro imperative verb prop	Ratio of imperative-like root verbs to all verb tokens in prompt introduction. Higher values indicate a more command- or instruction-oriented style.
H4 Specificity	title abstract token prop	Proportion of abstract tokens in prompt title, including lemmas in the preset abstract-word list (e.g., analysis, strategy, workflow) or lemmas ending in suffixes such as -tion, -ment, -ness, -ity, -ism, -ance, -ence, or -ship.
	title sentiment	VADER compound sentiment score for the prompt title. -1 = most negative, 0 = neutral, and +1 = most positive.
	intro concreteness	Concrete-token ratio minus abstract-token ratio in prompt introduction; concrete tokens are named entities, proper nouns, and non-abstract nouns; abstract tokens are lemmas in the preset abstract-word list (e.g., analysis, strategy, workflow) or lemmas ending in suffixes such as -tion, -ment, -ness, -ity, -ism, -ance, -ence, or -ship; higher values indicate more concrete wording.
H5 Creator	creator count	Number of creators contributing to prompt.
	creator prompt count	Total number of prompts created by creator.
Exploratory	title character count $\times$ title abstract token prop	
	intro concreteness $\times$ intro imperative verb prop	

Note. The table summarizes the predictors and interaction terms used in the main and robustness check models.

APPENDIX B. ROBUSTNESS CHECKS

Table B1 Robustness Check Model Results			
Hypothesis	Predictor	Platform-z Coefficient	Platform-z p
H1 Simplicity	title character count	0.0006	= .8246
	title avg word length	-0.0116	= .4855
	intro dependency distance	0.0089	= .8170
	intro readability	0.0078	< .0001
H2 Markup	title special token prop	1.0072	= .0191
	intro symbol prop	1.7675	= .0221
H3 Directiveness	title verb prop	-0.7159	= .0002
	intro imperative verb prop	-0.4566	< .0001
H4 Specificity	title abstract token prop	-0.4770	= .0130
	title sentiment	-0.0634	= .4939
	intro concreteness	-0.3754	= .0172
H5 Creator	creator count	0.0777	= .7711
	creator prompt count	0.0004	= .1280
Exploratory	title character count × title abstract token prop	-0.0123	= .3219
	intro concreteness × intro imperative verb prop	0.0518	= .9087

Note. Platform-z denotes the fixed-effects model estimated on the within-platform standardized $\ln(1 + \text{popularity})$ outcome.

Received: 06-July-2026, Manuscript No. AMSJ-26-17332; **Editor assigned:** 08-July-2026, PreQC No. AMSJ-26-17332(PQ); **Reviewed:** 22-July-2026, QC No. AMSJ-26-17332; **Revised:** 29-July-2026, Manuscript No. AMSJ-26-17332(R); **Published:** 06-Aug-2026