

Online Appendix A Platform AI Policies

This appendix documents the platform AI policies that ground our stance construct: the first subsection situates the two focal policy decisions within the broader cross-platform policy landscape, and the second reconstructs the timeline of AI-related actions on the two focal platforms.

A.1 Cross-Platform AI Policy Landscape

Tables A.1 and A.2 provide context for our two focal cases. Table A.1 situates the focal policies of Lofter and Graffiti Kingdom within the broader spectrum of AI-policy stance signals. Table A.2 documents the prevalence of the two specific intervention types we study (platform-developed AI tools and explicit AI prohibitions) across other comparable visual-art platforms.

Table A.1 Platform Policies and Policy-Communicated AI Stances: A Spectrum from Pro-AI to Anti-AI

Platform	Policy	Date	Stance (Pro-AI → Anti-AI)
Lofter	Developed a built-in Generative AI image generator.	2023-03	Active Embrace
Shutterstock	Contributor Fund compensating artists for data licensing; expanded six-year OpenAI partnership.	2023-07	Promotion / Partnership
CGTrader	No explicit policy; mostly silent on AI.	n/a	Permissive (Silent)
DeviantArt	Requires an explicit “Created with AI tools” label (auto-applied when metadata indicates AI).	2023-06	Permissive (Labeled)
Pixiv	Upload toggle for “AI-generated,” user search filters, and separate rankings for AI works.	2022-10	Permissive (Segregated)
Cubebrush	AI content permitted only if labeled, but then hidden from search and storefronts—“essentially invisible.”	2025-05	Restrictive (Labeled)
Dribbble	Community Guidelines: Do not share content that is wholly generated by AI. When creating content with AI assistance, it is recommended to acknowledge the contribution of AI in your Shot description.	n/a	Restrictive (Human-Primary)
Graffiti Kingdom	Community rules prohibit uploading AI-generated images. Any users who violate this will be punished.	2023-02	Active Restriction

Note. The unit of analysis is the platform-policy case. Policies are ordered from most pro-AI to most anti-AI, with the focal Lofter and Graffiti Kingdom policy decisions anchoring the opposite ends of the spectrum.

Table A.2 Platforms with Generative AI Policies Comparable to Lofter and Graffiti Kingdom

Platform	Region	Core Focus	AI-Policy Bucket	Key Policy Summary
Lofter	China	Broad creative community with strong visual-art, illustration, photography, and fandom subcommunities	Platform-developed AI feature	LOFTER introduced and tested an in-platform AI image generator/profile-picture tool (often described as a “Profile Picture Generator” or AI drawing feature), later narrowing or removing access after creator backlash.
DeviantArt	Global (U.S.-based)	Online art community and portfolio platform for digital/traditional artists	Platform-developed AI feature	DeviantArt operates DreamUp, its own text-to-image service and policy framework for AI image generation on-platform.
Picsart	Global	Creator platform for image/video editing and design with a large maker community	Platform-developed AI feature	Picsart is primarily aimed at human creators and integrates in-platform generative-AI tools, including an AI image generator and prompt-based editing features, into its creator workflow.
Clip Studio Paint	Japan / Global	Digital drawing and comics platform used by illustrators, manga artists, and studios	Platform-developed AI feature	Clip Studio Paint officially announced an experimental in-app ‘Image Generator palette’ based on Stable Diffusion for creators, and then withdrew it days later after backlash from artists.
Graffiti Kingdom	China	Original illustration / drawing community and artist platform	Explicit prohibition	Platform content rules state that AI-generated images may not be uploaded; repeat violators can have works removed and uploading privileges restricted.
Cara	Global	Artist social and portfolio platform	Explicit prohibition	AI-generated media are not allowed in portfolios on Cara; AI content may appear only for news/reporting/discussion outside portfolio content and must be accurately labeled.
Newgrounds Art Portal	Global (U.S.-based)	Creator community with dedicated portals for art, games, audio, and animation	Explicit prohibition	Newgrounds’ Art Portal does not allow AI-generated art, aside from limited disclosed ancillary use cases (for example, an AI-generated background accompanying otherwise human-made character art).
INPRNT	Global (U.S.-based)	Curated art print marketplace and artist storefront platform	Explicit prohibition	INPRNT’s content guidelines prohibit artworks generated completely via an automated AI or machine-learning process.
Fur Affinity	Global (U.S.-based)	Online art community and portfolio platform focused on furry art	Explicit prohibition	Fur Affinity’s upload policy states that content made using AI, machine learning, or similar generators is not allowed, including derivative works created from it.

Note. The unit of analysis is the platform-policy case. The table summarizes comparable visual-art or creator platforms with either platform-developed Generative AI features or explicit prohibitions of AI-generated content.

A.2 Timeline of AI-Related Actions on the Focal Platforms

To provide a fuller account of the two policy episodes, we searched official platform pages, contemporaneous Chinese-language news and legal coverage, and the focal announcements retained by our author team. We included platform-level product actions, content rules, public explanations, and implementation measures that changed or explained how AI creation or AI-generated content would be governed; user commentary and third-party discussion were used only as corroborating evidence. Table A.3 reports the resulting chronology, covering the focal decisions and the directly related clarifications, restrictions, withdrawal, remediation, and implementation actions.

Table A.3 Timeline of Publicly Documented AI-Related Actions and Announcements on Lofter and Graffiti Kingdom, July 2022–June 2023

Platform	Date	Publicly Documented Action or Announcement	Role in the Focal Policy Episode	Evidence
Lofter	March 6, 2023	Publicly opened the “Profile Picture Generator”; the limited-test version had been called the “Lofter Drawing Machine.” Users could enter text prompts to generate images. In statements issued that day, Lofter characterized the tool as an avatar-use feature, stated that it used open-source training data rather than Lofter user works, prohibited commercial use, and promised adjustments.	<i>Focal launch and immediate explanation.</i> The product launch constitutes the treatment. The same-day statements explained the intended use of the tool without reversing the launch.	Jiemian; Sohu/Bianews; Pandaily
Lofter	March 7, 2023	Announced that the feature entry would be moved to the Avatar Frame Center and that generated images would be limited to profile-picture use, without download or publication functions. Lofter also announced prospective governance measures, including a rule against presenting AI-generated content as original work, a related feedback channel, anti-AI-scraping capability, and product mechanisms to distinguish AI content from original work.	<i>Reactive scope restriction and prospective remediation.</i> The restrictions narrowed the use of the focal tool. The additional governance measures were announced as future responses to creator concerns rather than as separately initiated AI policies.	Archived official statement; JWView; Superpixel (archived)
Lofter	March 8, 2023	Lofter subsequently stated that it had taken the disputed image-generation feature offline.	<i>Withdrawal.</i> The action removed the focal product feature and restored the pre-launch product state; it did not establish a new affirmative AI-governance policy.	Blue Whale News; China Economic Net
Lofter	March 10, 2023	Published an in-app apology acknowledging creator dissatisfaction and announced a forthcoming “Creator Protection Plan.” The plan outlined prospective measures, including capabilities intended to prevent AI-related misappropriation and malicious scraping, a dedicated infringement-reporting channel, restrictions on presenting AI-generated content as original work, stronger AI-content identification and feedback handling, a planned separate section for AI-generated content, and a future community-partner mechanism. ^a	<i>Apology and prospective remedial program.</i> The plan was presented as a response to the launch controversy. This row records the measures announced by Lofter and does not treat the listed measures as having been implemented at that time.	Sina; 21st Century Business Herald
Lofter	March 16, 2023	Issued a second public apology, reiterated that the disputed feature had been taken offline, and again stated that user works had not been used for AI training. Lofter reported that the rule against presenting AI content as original was in force, that an anti-misappropriation and anti-scraping system and a dedicated infringement-reporting channel were active, and that 1,148 infringement complaints had been processed by March 15. ^b It continued to describe the “Creator Protection Plan” as a forthcoming initiative. ^a	<i>Second apology and platform-reported follow-up.</i> The statement reported progress on selected elements of the remedial program within the same policy episode; it did not document implementation of every element announced on March 10.	Blue Whale News; Securities Daily/Tencent
Graffiti Kingdom	February 22, 2023	Announced that AI-generated artwork was prohibited, that previously uploaded works would be reviewed, and that violations would trigger sanctions. Platform staff subsequently clarified that a sanctioned user’s account and previously uploaded non-AI works would be preserved.	<i>Focal prohibition.</i> The policy explicitly excluded AI-generated artwork and constitutes the treatment analyzed in the paper.	Announcement retained by the authors; author correspondence; official content rules; contemporaneous legal commentary
Graffiti Kingdom	February 23–June 30, 2023	Maintained and implemented the February 22 prohibition through review and sanctions for AI-generated uploads. Our search identified no separately initiated AI-specific product launch or policy change during the remaining window.	<i>Implementation and maintenance.</i> These actions operationalized the same prohibition rather than creating a separate contemporaneous AI-policy shock.	Author correspondence and platform records; official content rules

Note. The primary event-search window is July 1, 2022–June 30, 2023. Searches combined each platform’s Chinese and English names with terms including AI, AIGC, AI drawing, AI-generated content, avatar generator, creator protection, anti-scraping, prohibition, and content rules. The table reports the platform-level actions and announcements identified in the public record and the focal materials retained by the authors; it does not rule out unobserved internal deliberations or actions that were not publicly documented. The original February 22 Graffiti Kingdom announcement is retained by the authors but is no longer publicly indexed; the current official rules and contemporaneous March 2023 commentary provide external corroboration. We found no additional independently initiated AI-specific platform action before the focal decisions or after the listed follow-up actions within the primary search window. Lofter’s March 6–16 actions are presented separately to make the sequence visible, although the post-launch restrictions, withdrawal, apologies, and remedial announcements were reactive components of the same policy episode. All hyperlinks in this table were last accessed on July 22, 2026.

^a For the prospective Lofter measures, we conducted an additional implementation search covering the year through March 16, 2024. We found no subsequent public announcement or independently verifiable evidence confirming that the planned separate AI-content section or the Lofter Partner Plan was launched. The absence of publicly available evidence does not rule out unpublished, partial, or limited implementation.

^b The March 16 implementation claims were made by Lofter and reported in contemporaneous news coverage. We found no independent technical audit of the anti-misappropriation or anti-scraping system. Moreover, the reported 1,148 complaints were described as infringement complaints generally, rather than specifically as AI-related complaints.

Online Appendix B Data Collection Details

Lofter

Lofter hosts various content formats (painting, writing, photography, filmmaking). We focus on visual artists because the AI image generator primarily affects this group. We compiled a list of visual-arts tags (e.g., *drawing*, *painting*). Among semantically similar tags, *painting* is most popular and provides exhaustive uploads for arbitrary time intervals. We retrieved all works tagged *painting* from January 2020 to December 2023, then visited each artist’s homepage to collect their complete work history.

Graffiti Kingdom

Graffiti Kingdom maintains a “New Works Today” section indexing all daily uploads. We systematically scraped publicly reachable historical pages (approximately 40,000 pages), achieving near-complete coverage from 2019 onward. For each upload, we visited the artist’s homepage and extracted all their previously uploaded works.

Full-Sample Descriptive Statistics

Table B.1 reports descriptive statistics for the full creator-month panel on each platform.

Table B.1 Full-Sample Descriptive Statistics of Key Variables

Variable	Definition	N	Mean	Std.Dev.	Min	Max
<i>Panel A: Lofter</i>						
Num of Works	Number of works uploaded during the month	7,510,542	1.84	18.35	0	14,726
Active	Equals 1 if creator uploads ≥ 1 work	7,510,542	0.05	0.22	0	1
review_avg	Average number of reviews received	7,510,542	0.98	9.41	0	8,141
like_avg	Average number of likes received	7,510,542	68.87	1,085.66	0	534,697
<i>Panel B: Graffiti Kingdom</i>						
Num of Works	Number of works uploaded during the month	1,822,186	1.01	4.64	0	1,464
Active	Equals 1 if creator uploads ≥ 1 work	1,822,186	0.16	0.37	0	1
review_avg	Average number of reviews received	1,822,186	0.07	0.85	0	167
like_avg	Average number of likes received	1,822,186	1.77	10.62	0	1,629

Note. The unit of analysis is the creator-month. The full sample covers Lofter from January 2020 to December 2023 and Graffiti Kingdom from January 2019 to December 2023; *N* denotes the number of creator-month observations.

Online Appendix C TCI Cohort Construction and Validation Tests

Tables C.1 and C.2 report, for each treated cohort on Lofter and Graffiti Kingdom, the matched control cohorts, the cohort sample sizes, and the per-pair Kolmogorov–Smirnov and Granger causality test results referenced in Section 4.

Table C.1 Validation Results for Temporal Causal Inference: Kolmogorov–Smirnov and Granger
Causality Tests for Lofter

Treated cohort	Control cohorts	Treated creators	Control creators	Granger F-statistic	Granger p-value	Kolmogorov–Smirnov D-statistic	Kolmogorov–Smirnov p-value
8	11,12,13,14,15	6,512	28,973	1.9011	0.2400	0.1250	1.0000
9	12,13,14,15,16	5,647	27,981	3.1156	0.1378	0.1111	1.0000
10	13,14,15,16,17	5,238	25,976	8.3616	0.0341	0.1000	1.0000
11	14,15,16,17,18	5,200	24,174	4.3150	0.0924	0.0909	1.0000
12	15,16,17,18,19	6,121	23,034	1.0064	0.3618	0.0833	1.0000
13	16,17,18,19,20	5,742	21,859	1.3114	0.3040	0.0769	1.0000
14	17,18,19,20,21	5,882	23,493	0.0252	0.8802	0.0714	1.0000
15	18,19,20,21,22	6,028	23,196	1.0955	0.3432	0.0667	1.0000
16	19,20,21,22,23	4,208	21,971	56.0278	0.0007	0.1250	0.9999
17	20,21,22,23,24	4,116	19,586	2.2838	0.1911	0.1176	0.9999
18	21,22,23,24,25	3,940	17,129	1.6858	0.2508	0.1111	1.0000
19	22,23,24,25,26	4,742	13,833	0.5877	0.4779	0.0526	1.0000
20	23,24,25,26,27	4,853	13,319	0.0348	0.8593	0.0500	1.0000
21	24,25,26,27,28	5,842	13,812	0.5614	0.4874	0.0476	1.0000
22	25,26,27,28,29	3,819	13,807	11.4692	0.0195	0.0455	1.0000
23	26,27,28,29,30	2,715	13,724	4.6711	0.0831	0.0435	1.0000
24	27,28,29,30,31	2,357	13,311	0.7441	0.4278	0.0417	1.0000
25	28,29,30,31,32	2,396	12,236	0.8257	0.4052	0.0400	1.0000
26	29,30,31,32,33	2,546	10,617	0.3923	0.5586	0.0385	1.0000
27	30,31,32,33,34	3,305	10,771	0.0156	0.9056	0.0370	1.0000
28	31,32,33,34,35	3,208	10,906	0.4819	0.5185	0.0357	1.0000
29	32,33,34,35,36	2,352	11,214	1.0702	0.3483	0.0690	1.0000
30	33,34,35,36,37	2,313	12,484	0.1336	0.7297	0.0667	1.0000
31	34,35,36,37,38	2,133	15,505	0.5147	0.5052	0.0645	1.0000
32	35,36,37,38,39	2,230	18,748	0.7012	0.4406	0.0625	1.0000
33	36,37,38,39,40	1,589	22,013	0.6622	0.4528	0.0909	0.9995
34	37,38,39,40,41	2,506	23,081	2.6588	0.1639	0.0882	0.9996
35	38,39,40,41,42	2,448	21,329	3.9318	0.1042	0.0857	0.9997
36	39,40,41,42,43	2,441	18,043	0.7938	0.4138	0.0556	1.0000
37	40,41,42,43,44	3,500	13,661	0.8702	0.3937	0.0541	1.0000
38	41,42,43,44,45	4,610	9,409	0.0128	0.9142	0.0526	1.0000
39	42,43,44,45,46	5,749	7,260	1.4256	0.2860	0.0513	1.0000
40	43,44,45,46,47	5,713	7,790	0.4799	0.5193	0.0750	0.9999
41	44,45,46,47,48	3,509	8,218	0.0346	0.8597	0.0976	0.9914
42	45,46,47,48,49	1,748	8,137	1.3046	0.3051	0.0714	1.0000
43	46,47,48,49,50	1,324	7,694	8.3095	0.0345	0.0698	1.0000
44	47,48,49,50,51	1,367	7,372	0.1346	0.7287	0.0682	1.0000
45	48,49,50,51,52	1,461	6,132	0.1158	0.7474	0.0444	1.0000
46	49,50,51,52,53	1,360	5,619	0.2759	0.6218	0.0652	1.0000
47	50,51,52,53,54	2,278	5,652	0.5760	0.4821	0.0638	1.0000
48	51,52,53,54,55	1,752	5,546	0.2187	0.6597	0.0625	1.0000
49	52,53,54,55,56	1,286	5,298	5.1326	0.0728	0.0612	1.0000
50	53,54,55,56,57	1,018	5,105	3.0342	0.1420	0.0600	1.0000
51	54,55,56,57,58	1,038	4,693	3.0267	0.1424	0.0392	1.0000
52	55,56,57,58,59	1,038	4,252	0.0036	0.9543	0.0385	1.0000
53	56,57,58,59,60	1,239	4,646	0.0803	0.7883	0.0566	1.0000
54	57,58,59,60,61	1,319	4,921	0.0607	0.8152	0.0370	1.0000
55	58,59,60,61,62	912	4,866	2.0041	0.2160	0.0545	1.0000

Note. The table-level unit of analysis is a treated-cohort/matched-control-cohort pair; the underlying data are creator-month observations. Cohort i refers to creators whose first work was posted i months before Lofter’s launch of the “Profile Picture Generator”. Rows with Granger p-values below 0.05 fail the validation screen and are excluded from the subsequent Temporal Causal Inference estimation; all listed pairs pass the Kolmogorov–Smirnov test.

Table C.2 Validation Results for Temporal Causal Inference: Kolmogorov–Smirnov and Granger
Causality Tests for Graffiti Kingdom

Treated cohort	Control cohorts	# Treated creators	# Control creators	Granger F-statistic	Granger p-value	Kolmogorov–Smirnov D-statistic	Kolmogorov–Smirnov p-value
7	10,11,12,13,14	925	5,582	0.0008	0.9791	0.1429	1.0000
8	11,12,13,14,15	1,041	5,951	0.0688	0.8101	0.1250	1.0000
9	12,13,14,15,16	1,200	5,848	0.1981	0.6864	0.1111	1.0000
10	13,14,15,16,17	963	5,720	11.7603	0.0416	0.1000	1.0000
11	14,15,16,17,18	1,091	5,656	4.3461	0.1284	0.0909	1.0000
12	15,16,17,18,19	1,200	5,582	0.8339	0.4285	0.0833	1.0000
13	16,17,18,19,20	1,188	5,260	1.0625	0.3785	0.1538	0.9992
14	17,18,19,20,21	1,140	4,966	2.9061	0.1868	0.1429	0.9996
15	18,19,20,21,22	1,332	4,498	5.5025	0.1007	0.1333	0.9998
16	19,20,21,22,23	988	3,892	3.2195	0.1706	0.0625	1.0000
17	20,21,22,23,24	1,072	3,367	0.8487	0.4249	0.1176	0.9999
18	21,22,23,24,25	1,124	2,862	0.7677	0.4454	0.1111	1.0000
19	22,23,24,25,26	1,066	2,717	34.5998	0.0098	0.1579	0.9781
20	23,24,25,26,27	1,010	2,488	1.3147	0.3347	0.1000	1.0000
21	24,25,26,27,28	694	2,428	9.0789	0.0571	0.0952	1.0000
22	25,26,27,28,29	604	2,283	4.3433	0.1285	0.0455	1.0000
23	26,27,28,29,30	518	2,162	4.4171	0.1264	0.0870	1.0000
24	27,28,29,30,31	541	1,986	12.4283	0.0388	0.0417	1.0000
25	28,29,30,31,32	505	1,995	6.1046	0.0900	0.0800	1.0000
26	29,30,31,32,33	549	2,086	4.8862	0.1141	0.0385	1.0000
27	30,31,32,33,34	375	2,455	1.4238	0.3185	0.0370	1.0000
28	31,32,33,34,35	458	2,902	2.3572	0.2223	0.0714	1.0000
29	32,33,34,35,36	396	3,376	0.6368	0.4832	0.0690	1.0000
30	33,34,35,36,37	384	4,113	2.8048	0.1926	0.1000	0.9988
31	34,35,36,37,38	373	4,542	0.0288	0.8759	0.1290	0.9634
32	35,36,37,38,39	384	4,564	29.3603	0.0123	0.0625	1.0000
33	36,37,38,39,40	549	4,112	7.4284	0.0722	0.0606	1.0000
34	37,38,39,40,41	765	3,515	11.2943	0.0437	0.0588	1.0000
35	38,39,40,41,42	831	2,663	1.1378	0.3643	0.0571	1.0000
36	39,40,41,42,43	847	1,968	0.4950	0.5324	0.0833	0.9998
37	40,41,42,43,44	1,121	1,481	0.6806	0.4699	0.1081	0.9846
38	41,42,43,44,45	978	1,457	9.5919	0.0534	0.1579	0.7379
39	42,43,44,45,46	787	1,474	3.1554	0.1738	0.1282	0.9114
40	43,44,45,46,47	379	1,518	0.7627	0.4468	0.1250	0.9188
41	44,45,46,47,48	250	1,578	0.0155	0.9088	0.0976	0.9914
42	45,46,47,48,49	269	1,546	0.3189	0.6117	0.0952	0.9926

Note. The table-level unit of analysis is a treated-cohort/matched-control-cohort pair; the underlying data are creator-month observations. Cohort i refers to creators whose first work was posted i months before Graffiti Kingdom’s prohibition of AI-generated artwork. Rows with Granger p-values below 0.05 fail the validation screen and are excluded from the subsequent Temporal Causal Inference estimation; all listed pairs pass the Kolmogorov–Smirnov test.

Online Appendix D Identification Assumptions & Validation

Before presenting the estimation results, we discuss the identification assumptions following the Potential Outcome Framework (Rubin 1980) to validate the TCI framework in our context.

A1: SUTVA (Stable Unit Treatment Value Assumption)

The SUTVA assumption (Rubin 1980) consists of two conditions. First, there should be no interference between units. In our case, this means that whether one content creator is exposed to the treatment should not affect the outcomes of any other content creators on the platform. In other words, $Y_i^{(W_i)} \perp Y_j^{(W_j)}$ for any two content creators $i, j \in 1, 2, \dots, N$. Under our TCI framework, the assignment to treatment or control groups is primarily based on the cohort to which the content creator belongs. This implies that the outcomes do not occur at the same real-world time (Turjeman and Feinberg 2023), meaning that no two users in different groups (i.e., control and treatment) can influence each other's treatment or lack thereof (i.e., $\forall i, j, Y_i^{(1)} \perp Y_j^{(0)}$).

An important concern here is the network effect. First, due to the potential network effect, $Y_i^{(1)} \perp Y_j^{(1)}$ and $Y_i^{(0)} \perp Y_j^{(0)}$ may not hold for some i, j pairs. For instance, users may increase or decrease their activity in response to the behavior of other users they know. However, since our objective is to measure how the AI stances signaled by platform policy decisions affect users' behavior, this potential dependency is a crucial aspect of estimating the treatment effect and should be considered in its assessment (Turjeman and Feinberg 2023).

Second, the potential network effect means that, as earlier users joining the platform, the activities of content creators in the control group may influence those in the treatment group. However, both platforms are stable and mature as they have been established for over ten years. Moreover, the treatment assignment is independent of this potential effect since the treatments (i.e., the platform policies) could not be anticipated. Even if interference does exist between units – where the behavior of content creators in the control group,

who joined before the treatment group, influences those in the treatment group – a similar argument applies to the control group concerning their predecessors. This means the activities of content creators in the control group are also affected by those who joined earlier. Therefore, any similar effects, if they do exist, would be accounted for by simply comparing the control and treatment units (Turjeman and Feinberg 2023). Moreover, we apply the Granger causality test (Turjeman and Feinberg 2023, Granger 1980): For each treatment–control pair, the test evaluates whether behavior of the control cohort (earlier entrants) has predictive power for behavior in the treatment cohort. Evidence of predictive power implies a breach of SUTVA; accordingly, we classify that treatment–control pair as failing the Granger test, which are not used in our analysis.

Additionally, SUTVA requires that no hidden versions of treatments exist. In other words, nothing except the treatment itself can influence the outcome variables, i.e., $Y_i = Y_i^{(1)}W_i + Y_i^{(0)}(1 - W_i)$. We meticulously reviewed all previous announcements from the two platforms and confirmed that, at the time of the treatments, there were no other changes to the platforms, significant holidays, or confounding events apart from the treatments themselves. We also examined other prominent digital art platforms in China to verify that they did not take any actions related to Generative AI during the same period. In particular, we check the announcements of the following platforms: ZCOOL, Mihuashi, bcy.net, huaban.com, and duitang.com, all of which are well-known Chinese platforms and can be considered competitors to Lofter and Graffiti Kingdom. Moreover, our TCI framework, which incorporates multiple cohorts into the control group, helps mitigate the impact of unobserved events on the outcome variable Y_i . Any potential effects from such events can be averaged out across all control cohorts, providing a more reliable control group baseline (Turjeman and Feinberg 2023). Furthermore, we apply the Kolmogorov-Smirnov (KS) test to demonstrate that there are no significant differences in behaviors between the control groups constructed by TCI and the treatment groups prior to the treatments. This is discussed in more detail in the next assumption.

A2: Conditional Independence Assumption

This assumption states that, the assignment of treatment or control group should be (approximately) random, conditional on the control variables X_i , ensuring the comparability between treatment and control units:

$$Pr(W_i|X_i, Y^{(0)}, Y^{(1)}) = Pr(W_i|X_i)$$

While the treatments on both platforms can be considered exogenous since the policy shocks are unanticipated, one potential concern is that the joining time may be confounded with user activities. Under the TCI framework, the constructed control groups consist of content creators who joined before those in the treatment groups. Therefore, if unobserved factors influence both the joining time and user activities on the platforms, our estimated treatment effects would be biased.

To address this concern, we apply the Kolmogorov-Smirnov (KS) Test to determine whether the control and treatment groups differ in their activities at the same tenures (Dixon and Massey Jr 1951, Turjeman and Feinberg 2023). Specifically, for each treatment-control pair, we calculate the cumulative sum of the average number of works uploaded to the platforms before the treatment. We then divide this by the maximum cumulative sum of the groups, which create a CDF-like timeline that can be tested using the KS Test (Turjeman and Feinberg 2023).

Only treatment-control pairs that pass both KS and Granger tests are included in our subsequent analysis. Columns (5)-(8) of Tables C.1 and C.2 show the KS Test and Granger Causality Test's results for Lofter and Graffiti Kingdom, separately. We can see, all the cohorts (on both platforms) have passed the KS test while a few cohorts do not pass the Granger test, and hence will not be used in our later analysis. Since only a small number of cohorts fail the Granger test and therefore excluded from the analysis, the impact on our sample size is minimal.

Last but not least, to test the parallel trends assumption for the DiD specification, we run an event study for each cohort and aggregate the resulting coefficients across cohorts

using the same inverse-cohort-size weighting described above. We normalize the period immediately before treatment as the base period and interpret all event-time coefficients relative to this pre-treatment benchmark. As seen in Figure 3 and its corresponding Table D.1, the pre-period coefficients do exhibit some mild directional drift visually, but they are not statistically different from zero, indicating no meaningful pre-trends.¹

A3: Overlap Assumption

This assumption requires that, for each content creator, her propensity to be treated should be strictly between 0 and 1:

$$0 < P(W_i = 1 | X_i = x) < 1$$

Due to the nature of the exogenous treatments, every content creator who joined the platforms before the treatment has a positive chance of being treated. In fact, all content creators in our main sample were eventually treated, but at different tenures on the platforms.

Finally, Table 1 reports summary statistics for the final estimation sample. This sample retains only observations from cohorts that pass both the Granger causality test and the KS test. In addition, as mentioned earlier, we exclude users who joined more than 55 months earlier before the focal date on Lofter and more than 42 months on Graffiti Kingdom, so as to restrict attention to cohorts that are more comparable in their pre-treatment behavioral dynamics. We note that the reduction in sample size is somewhat larger for Graffiti Kingdom than for Lofter. One reason is that Graffiti Kingdom has relatively earlier-entry cohorts. Because these cohorts correspond to users who had been on the platform longer, they contribute more observations to the panel. As a result, excluding them removes a larger share of the final sample.

¹The wider confidence intervals for Lofter in later event-time periods do not reflect compositional change. In our setting, users who become inactive remain in the sample, with their number of works recorded as zero, so the estimates are not driven by attrition of inactive users. Rather, the widening confidence intervals primarily reflect greater heterogeneity in creators' post-treatment behavior.

Table D.1 Parallel-Trends Test under Temporal Causal Inference

	Lofter	Graffiti Kingdom
	(1)	(2)
Dependent Variable	log(Num of Works + 1)	
Treat $\times$ Pre Month 7	-0.0033 [0.0076]	0.0045 [0.0095]
Treat $\times$ Pre Month 6	-0.0085 [0.0057]	0.0035 [0.0088]
Treat $\times$ Pre Month 5	-0.0106 [0.0067]	0.0041 [0.0088]
Treat $\times$ Pre Month 4	-0.0094 [0.0065]	0.0042 [0.0079]
Treat $\times$ Pre Month 3	-0.0048 [0.0047]	-0.0012 [0.0079]
Treat $\times$ Pre Month 2	0.0007 [0.0041]	0.0019 [0.0072]
Treat $\times$ Pre Month 1	0 [—]	0 [—]
Treat $\times$ Post Month 1	-0.0894*** [0.0077]	0.0169** [0.0074]
Treat $\times$ Post Month 2	-0.1707*** [0.0079]	0.0146* [0.0081]
Treat $\times$ Post Month 3	-0.1707*** [0.0081]	0.0055 [0.0091]
Creator Fixed Effects	Yes	Yes
Tenure Fixed Effects	Yes	Yes
Number of Creators	140,061	24,178

Note. The unit of analysis is the creator-month; Pre Month 1 is the omitted reference period and its coefficient is normalized to zero. Robust standard errors clustered at the creator level are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Online Appendix E Hazard Regression & Results

To assess whether policy timings are strategically aligned with creators' activities on the platforms, we estimate a discrete-time hazard model on platform-month aggregates prior to treatments ($t \leq 1$, with $t = 1$ denoting the happening months of the treatments). Let Posts_{pt} (Num of Works), Likes_{pt} (Num of Likes), and Reviews_{pt} (Num of Reviews) be monthly totals. Here, the number of likes and the number of reviews denote the total counts received by works uploaded in month t . Because we do not observe time-varying data on

the likes and reviews that each work receives, we use these variables as proxies for the actual number of likes and the actual number of reviews on the whole platform during month t . We construct lagged levels $\text{lvl_}x_{pt} = x_{p,t-1}$, six-month rolling trends $\text{slp_}x_{pt}$ (OLS slope within the past window), and volatilities $\text{vol_}x_{pt}$ (within-window standard deviation), for $x \in \{\text{Posts, Likes, Reviews}\}$. The adoption indicator is defined as

$$\text{adopt}_{pt} = \mathbb{I}\{t = 1\},$$

which equals 1 in the month when the treatment occurs and 0 otherwise.

We then estimate a hazard regression to assess whether pre-event dynamics predict adoption timing, with the model specification as follows:

$$\Pr(\text{adopt}_{pt} = 1 \mid \mathcal{H}_{p,t-1}) = \text{logit}^{-1}(Z_{pt}),$$

where $Z_{pt} = \beta_0 + \beta_1 \text{lvl_posts}_{pt} + \beta_2 \text{slp_posts}_{pt} + \beta_3 \text{vol_posts}_{pt} + \beta_4 \text{lvl_likes}_{pt} + \beta_5 \text{slp_likes}_{pt} + \beta_6 \text{vol_likes}_{pt} + \beta_7 \text{lvl_reviews}_{pt} + \beta_8 \text{slp_reviews}_{pt} + \beta_9 \text{vol_reviews}_{pt} + \gamma t$. Table E.1 reports the estimates. None of these coefficients is statistically significant, providing no evidence that pre-event platform-level dynamics predict policy adoption timing.

Table E.1 Discrete-Time Hazard Regression Results for Policy Adoption Timing

Dependent Variable	Lofter	Graffiti Kingdom
	(1)	(2)
	Adopt	
(Intercept)	7.6121 [16.6743]	-7.6985 [13.1650]
lvl_posts	-0.0001 [0.0001]	0.0001 [0.0005]
slp_posts	0.0001 [0.0002]	-0.0026 [0.0040]
vol_posts	0.0001 [0.0002]	0.0018 [0.0018]
lvl_likes	0.00001 [0.0001]	-0.0001 [0.0002]
slp_likes	-0.0001 [0.0001]	0.0007 [0.0010]
vol_likes	0.00001 [0.0001]	-0.0003 [0.0003]
lvl_reviews	0.0001 [0.0003]	0.0010 [0.0025]
slp_reviews	0.0001 [0.0004]	-0.0122 [0.0153]
vol_reviews	-0.0001 [0.0002]	0.0061 [0.0091]
time_trend	0.1389 [0.3396]	-0.0840 [0.2369]
Number of Observations	31	31

Note. The unit of analysis is the platform-month. The dependent variable equals one in the policy-adoption month and zero otherwise; each column uses 31 monthly observations ending in that month (Lofter: March 6, 2023; Graffiti Kingdom: February 22, 2023). The covariates are lagged levels (*lvl*), six-month rolling slopes (*slp*), and six-month rolling volatilities (*vol*) of monthly platform-level totals of works/posts, likes, and reviews. Robust standard errors clustered at the month level are in brackets. $*p < 0.1$, $**p < 0.05$, $***p < 0.01$.

Online Appendix F Long-Term Signaling Effect of Lofter's Pro-AI Policy Signal

Given the estimates in Table 2, we observe negative and significant effects for Lofter and positive and significant effects for Graffiti Kingdom in the first three post-treatment months. This naturally raises the question of whether these effects persist over a longer horizon.

Accordingly, we extend the post-treatment window to six months and re-estimate the TCI for both platforms. As reported in Table F.1, the estimated coefficients remain statistically significant in all six post-treatment months for both Lofter and Graffiti Kingdom. For Lofter, all coefficients remain negative, consistent with a persistent negative signaling effect of Lofter's pro-AI policy signal on creators' behavior. For Graffiti Kingdom, all coefficients remain positive, consistent with a positive effect that holds over the six-month window. It is worth noting that, because the post-treatment window now spans six months, we exclude control cohorts that entered less than six months before the treated cohorts, so that the control group remains uncontaminated over the full post-treatment horizon.

Table F.1 Six-Month Temporal Causal Inference Results Using Difference-in-Differences

Dependent Variable	Lofter	Graffiti Kingdom
	(1) log(Num of Works + 1)	(2)
Treat $\times$ Post Month 1	-0.0935*** [0.0075]	0.0283*** [0.0085]
Treat $\times$ Post Month 2	-0.1170*** [0.0065]	0.0274*** [0.0085]
Treat $\times$ Post Month 3	-0.0964*** [0.0067]	0.0202** [0.0084]
Treat $\times$ Post Month 4	-0.0906*** [0.0066]	0.0206** [0.0096]
Treat $\times$ Post Month 5	-0.0832*** [0.0067]	0.0251*** [0.0097]
Treat $\times$ Post Month 6	-0.0771*** [0.0053]	0.0315*** [0.0094]
Estimation Method	Difference-in-Differences	
Creator Fixed Effects	Yes	Yes
Tenure Fixed Effects	Yes	Yes
Number of Creators	140,061	27,897

Note. The unit of analysis is the creator-month; Post Month 1–6 denote the first six post-treatment periods. Robust standard errors clustered at the creator level are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Online Appendix G Robustness Checks

Our main results rely on the Temporal Causal Inference framework with Difference-in-Differences estimation and log-transformed outcome variables. To ensure robustness, we conduct several sets of robustness checks that address potential concerns about estimation method, identification strategy, outcome measurement, functional form, and sample selection. All robustness checks yield results consistent with our main findings, which strengthens confidence in our conclusions.

First, we apply Generalized Synthetic Control (GSC) as an alternative estimator within the TCI framework. Second, we implement a Regression Discontinuity in Time (RDiT) design (Cavusoglu et al. 2016, Hausman and Rapson 2018, Shi et al. 2022), which exploits the discontinuity at the policy cutoff rather than tenure-based control groups. Third, we use a Panel Difference Estimator (Goldberg et al. 2024), which compares year-over-year changes for the same creators. Fourth, motivated by concerns about log-like transformations $\log(Y + 1)$ (Chen and Roth 2024), we use an alternative binary outcome variable (“active” indicator). Last, we replace the linear specification with Poisson pseudo-maximum likelihood (PPML) estimation. These checks provide strong support for our main findings.

G.1 Generalized Synthetic Control

As an alternative to DiD, Generalized Synthetic Control (GSC) (Xu 2017) constructs a synthetic control group as a weighted combination of control units that best matches the treatment group’s pre-treatment trajectory. While DiD assumes parallel trends between treatment and control groups, GSC relaxes this assumption by allowing for interactive fixed effects. We apply GSC within the same TCI framework, using the same treatment–control pairs constructed by cohort matching.

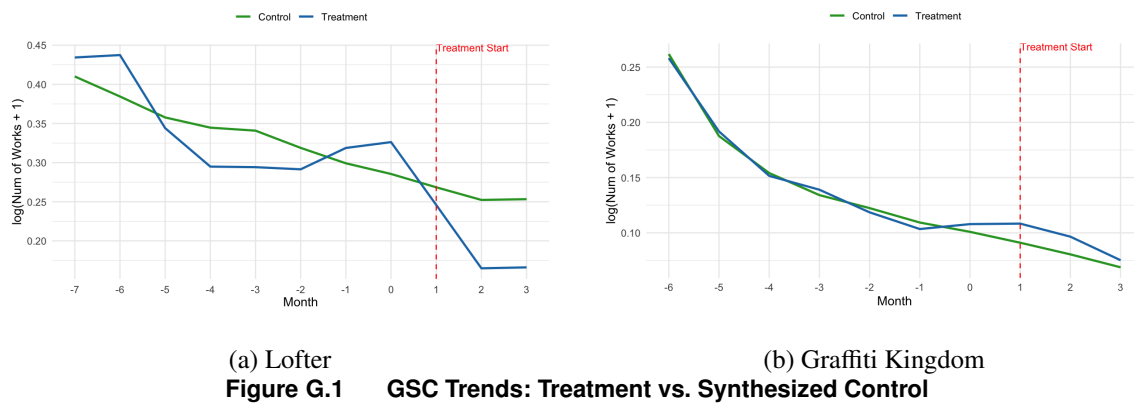
Table G.1 presents GSC estimates. The results are directionally consistent with our DiD estimates in Table 2: negative treatment effects for Lofter and positive effects for Graffiti Kingdom. Figure G.1 visualizes the treatment and synthesized control group trends. We note that the pre-treatment fit is imperfect. For Lofter, the treatment and control groups

show some divergence in pre-treatment levels; for Graffiti Kingdom, a small gap appears in period 0, the month before the policy was implemented. Despite these limitations, the directional consistency between GSC and DiD estimates provides reassurance for our main findings.

Table G.1 Temporal Causal Inference Results Using Generalized Synthetic Control

	Lofter	Graffiti Kingdom
	(1)	(2)
Dependent Variable	log(Num of Works + 1)	
Treat $\times$ Post Month 1	-0.0280*** [0.0065]	0.0173*** [0.0049]
Treat $\times$ Post Month 2	-0.0916*** [0.0058]	0.0160*** [0.0050]
Treat $\times$ Post Month 3	-0.0910*** [0.0056]	0.0064 [0.0043]
Estimation Method	Generalized Synthetic Control	
Creator Fixed Effects	Yes	Yes
Tenure Fixed Effects	Yes	Yes
Number of Creators	140,061	24,178

Note. The unit of analysis is the creator-month; Post Month 1–3 denote the first three post-treatment periods. Standard errors computed using nonparametric bootstrapping are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.



G.2 Regression Discontinuity in Time

TCI relies on tenure-based control groups and requires that creators at the same tenure exhibit comparable behavior across cohorts. RDiT offers an alternative identification strategy that does not require tenure matching. Instead, RDiT exploits the discontinuity at the policy implementation date and differences out unobservable factors that evolve smoothly over time.

We restrict the sample to three months before and three months after the policy change – a window sufficiently wide to observe effects yet narrow enough to enhance comparability.² Following Cavusoglu et al. (2016), we estimate:

$$\begin{aligned} Y_{it} = & \lambda_i + \beta_1 \cdot AfterPolicy_t + \beta_2 \cdot Month_t \\ & + \beta_3 \cdot AfterPolicy_t \times Month_t + \epsilon_{it}. \end{aligned} \tag{G.1}$$

where i indexes content creators and t indexes months (not tenure). The outcome variable Y_{it} is the log of the number of works uploaded by content creator i to the platform during month t . We code $Month_t$ as the number of months before or after the policy change, which serves as the running variable in the RDiT model (Cavusoglu et al. 2016, Hausman and Rapson 2018). $AfterPolicy_t = 1$ if month t falls within the post-treatment period, and 0 otherwise; λ_i represents the fixed effects of individual creator i , and ϵ_{it} is the error term. Our primary coefficient of interest, β_1 , captures the magnitude of the discontinuous change in the outcome variable at the cutoff month. The interaction term $AfterPolicy_t \times Month_t$ allows us to fit different time trends on each side of the cutoff point (Cavusoglu et al. 2016).

Table G.2 presents RDiT estimates. For Lofter, the discontinuous drop at the policy cutoff is -0.1144 ($p < 0.001$). For Graffiti Kingdom, the discontinuous jump is $+0.0103$ ($p < 0.001$). These estimates align closely with our TCI results in Table 2, which confirm negative signaling effects from Lofter’s pro-AI policy signal and positive signaling effects from Graffiti Kingdom’s anti-AI policy signal.

²We therefore exclude users who created their accounts less than three months before the policy change, i.e., those who had not joined the platform before the start of our three-month pre-policy window.

Table G.2 Regression Discontinuity in Time Results

	Lofter	Graffiti Kingdom
Dependent Variable	(1) log(Num of Works + 1)	(2) log(Num of Works + 1)
AfterPolicy	-0.1144*** [0.0020]	0.0103*** [0.0029]
Month	0.0128*** [0.0008]	-0.0050*** [0.0011]
AfterPolicy × Month	-0.0569*** [0.0011]	-0.0098*** [0.0016]
Creator Fixed Effects	Yes	Yes
Adj. R-squared	0.0258	0.0015
Number of Creators	186,660	37,576
Observation Window Length	6	6
Number of Observations	1,119,960	225,456

Note. The unit of analysis is the creator-month. The Regression Discontinuity in Time analysis uses a six-month window, with three monthly periods before and three monthly periods after each policy cutoff (Lofter: March 6, 2023; Graffiti Kingdom: February 22, 2023). Robust standard errors clustered at the creator level are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

G.3 Panel Difference Estimator

As a second alternative to TCI, the Panel Difference Estimator uses prior-year observations for the same creators as controls. This design absorbs seasonal variation (e.g., holiday effects) while remaining insulated from policy changes, which occurred only in 2023. The method is analogous to Difference-in-Differences and has been widely used in economics (Goldberg et al. 2024, Eichenbaum et al. 2024, Babina et al. 2024) and information systems (Burtch et al. 2023).³

We focus on creators observed in at least two consecutive years (2022 and 2023), using 2022 observations as controls and 2023 observations as treatment, aligned by calendar month. We estimate:

$$Y_{it} = \mu_i + \kappa_t + \sum_{m=1}^3 \tau_m \text{Treat}_i \times \text{PostMonth}_{mt} + \epsilon_{it}. \quad (\text{G.2})$$

³Also called the "long-difference estimator" in financial economics and macroeconomics.

where i indexes creators and t indexes calendar months; Y_{it} is the log of works uploaded; $Treat_i = 1$ for 2023 observations (zero for 2022); $PostMonth_{mt}$ is an indicator equal to one if period t corresponds to post-policy month $m \in \{1, 2, 3\}$ and zero otherwise; μ_i are individual fixed effects; κ_t are calendar month fixed effects; and ϵ_{it} is the error term. The coefficients τ_1 , τ_2 , and τ_3 capture the time-varying treatment effects in the three post-policy months.

Table G.3 presents estimates using both DiD and GSC. Results align closely with our TCI estimates in Table 2, again confirming negative signaling effects on Lofter and positive signaling effects on Graffiti Kingdom.

Table G.3 Panel Difference Estimator Results Using Difference-in-Differences and Generalized Synthetic Control

Dependent Variable	Lofter		Graffiti Kingdom	
	(1)	(2) log(Num of Works + 1)	(3)	(4)
Treat $\times$ Post Month 1	-0.0476** [0.0207]	-0.0482*** [0.0023]	0.0628** [0.0238]	0.0568*** [0.0035]
Treat $\times$ Post Month 2	-0.1048*** [0.0234]	-0.1320*** [0.0024]	0.0903*** [0.0295]	0.0568*** [0.0035]
Treat $\times$ Post Month 3	-0.0782*** [0.0234]	-0.1395*** [0.0025]	0.1476*** [0.0295]	0.0555*** [0.0043]
Estimation Method	Difference-in-Differences	Generalized Synthetic Control	Difference-in-Differences	Generalized Synthetic Control
Creator Fixed Effects	Yes	Yes	Yes	Yes
Calendar Month Fixed Effects	Yes	Yes	Yes	Yes
Adj. R-squared	0.4754		0.3762	
Number of Creators	156,841	156,841	39,799	39,799
Observation Window Length	26	26	26	26
Number of Observations	4,077,866	4,077,866	1,034,744	1,034,744

Note. The unit of analysis is the creator-month. The Panel Difference Estimator compares policy-year observations in 2023 with aligned calendar-month observations in 2022; Post Month 1–3 denote the first three post-policy months and their aligned 2022 comparison periods. For Difference-in-Differences, robust standard errors clustered at the creator level are in brackets; for Generalized Synthetic Control, bootstrapped standard errors are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

G.4 Alternative Outcome Measures and Functional Forms

Our main results use $\log(Y + 1)$ transformations. Chen and Roth (2024) show that if treatment operates on the extensive margin, re-scaling Y before the log transformation can yield treatment effects of any desired size. We address this concern using three approaches.

First, we conduct a scaling-sensitivity analysis using alternative offsets in the log transformation, replacing $\log(Y + 1)$ with $\log(Y + 2)$, $\log(Y + 3)$, and $\log(Y + 4)$. The corresponding estimates are reported in Table G.4. Reassuringly, the qualitative conclusions remain unchanged across all specifications: the estimated effects retain the same signs, remain statistically significant, and are broadly similar in magnitude to those reported in Table 2. These results suggest that our main findings are not driven by the particular choice of the +1 offset in the baseline log transformation.

Table G.4 Scaling-Sensitivity Analysis under Temporal Causal Inference Using Difference-in-Differences

Dependent Variable Log Offset	Lofter			Graffiti Kingdom		
	(1)	(2)	(3)	(4)	(5)	(6)
	log(Num of Works + offset)					
	+2	+3	+4	+2	+3	+4
Treat × Post Month 1	-0.0379*** [0.0046]	-0.0307*** [0.0037]	-0.0261*** [0.0032]	0.0192*** [0.0037]	0.0155*** [0.0030]	0.0131*** [0.0026]
Treat × Post Month 2	-0.0832*** [0.0043]	-0.0669*** [0.0035]	-0.0567*** [0.0030]	0.0184*** [0.0037]	0.0149*** [0.0030]	0.0127*** [0.0025]
Treat × Post Month 3	-0.0814*** [0.0042]	-0.0655*** [0.0034]	-0.0554*** [0.0029]	0.0108*** [0.0034]	0.0086*** [0.0028]	0.0073*** [0.0024]
Estimation Method	Difference-in-Differences					
Creator Fixed Effects	Yes	Yes	Yes	Yes	Yes	Yes
Tenure Fixed Effects	Yes	Yes	Yes	Yes	Yes	Yes
Number of Creators	140,061	140,061	140,061	24,178	24,178	24,178

Note. The unit of analysis is the creator-month. Post Month 1–3 denote the first three post-treatment periods. Robust standard errors clustered at the creator level are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Second, we replace the continuous outcome with a binary “active” indicator (= 1 if creator uploads at least one work in month t). We re-estimate Equation (1) under the TCI

framework with this alternative outcome. Table G.5 reports the results, which show patterns similar to Table 2.

Table G.5 Temporal Causal Inference Results Using Difference-in-Differences with Active

Dependent Variable	Outcome	
	Lofter	Graffiti Kingdom
	(1)	(2)
	Active	
Treat $\times$ Post Month 1	-0.0280*** [0.0034]	0.0156*** [0.0028]
Treat $\times$ Post Month 2	-0.0646*** [0.0031]	0.0143*** [0.0030]
Treat $\times$ Post Month 3	-0.0657*** [0.0029]	0.0086*** [0.0026]
Estimation Method	Difference-in-Differences	
Creator Fixed Effects	Yes	Yes
Tenure Fixed Effects	Yes	Yes
Number of Creators	140,061	24,178

Note. The unit of analysis is the creator-month; Post Month 1–3 denote the first three post-treatment periods. Robust standard errors clustered at the creator level are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Last but not least, following Chen and Roth (2024), we replace the linear model with Poisson regression estimated via Poisson pseudo-maximum likelihood (PPML). PPML does not require the full Poisson distributional assumption (e.g., $Var(Y | X) = \mathbb{E}[Y | X]$) (Gourieroux et al. 1984b,a) and is recommended as a substitute for log-like transformations (Chen and Roth 2024, Silva and Tenreiro 2006). The specification is:

$$\mathbb{E}[Y_{it} | \text{Treat}_i, \text{PostMonth}_{mt}, \mu_i, \lambda_t] = \exp\left(\mu_i + \lambda_t + \sum_{m=1}^3 \tau_{jm} \text{Treat}_i \times \text{PostMonth}_{mt}\right). \quad (\text{G.3})$$

with variable definitions identical to those in Equation (1). Table G.6 reports PPML estimates. Results broadly align with our main findings in Table 2, indicating that log-like transformation concerns do not drive our results.

Table G.6 Temporal Causal Inference Results Using Poisson Pseudo-Maximum Likelihood

	Lofter	Graffiti Kingdom
	(1)	(2)
Dependent Variable	Num of Works	
Treat $\times$ Post Month 1	-0.2618*** [0.0270]	0.1701** [0.0673]
Treat $\times$ Post Month 2	-0.6067*** [0.0322]	0.1235* [0.0690]
Treat $\times$ Post Month 3	-0.5839*** [0.0375]	-0.0736 [0.0715]
Estimation Method	Poisson Pseudo-Maximum Likelihood	
Creator Fixed Effects	Yes	Yes
Tenure Fixed Effects	Yes	Yes
Number of Creators	140,061	24,178

Note. The unit of analysis is the creator-month; Post Month 1–3 denote the first three post-treatment periods. Robust standard errors clustered at the creator level are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

In summary, all robustness checks – alternative estimation methods (GSC), alternative identification strategies (RDiT, PDE), alternative outcomes (binary active indicator), and alternative functional forms (PPML) – yield results broadly consistent with our main findings.

G.5 Sample Selection

Finally, we conduct four additional sample-selection and comparison-group robustness checks, and report the corresponding regression results in Table G.7, which are qualitatively consistent with our main estimation results. First, to address the concern that our results may be driven by a small number of highly productive users, we re-estimate the baseline specification after excluding *super-users*, defined as the top 1% of users in each cohort ranked by total work volume (`work_sum`). We chose this cutoff because the data cover the entire user population, most of whom post little or nothing in a given month. With so much of the distribution concentrated at the bottom, even the top 1% threshold reaches down to a modest 24 works a month on Lofter and 18 on Graffiti Kingdom, so a wider cut would remove ordinary active creators rather than additional super-users. Reassuringly, excluding

the top 1% of users leaves our main results qualitatively unchanged, indicating that the estimated treatment effects are not driven by a small group of dominant contributors.

Second, we examine whether our findings are sensitive to the precise choice of nearby control cohorts. In the baseline design, for each treated cohort i , the control group is constructed from cohorts $i + 3$ to $i + 7$. As a robustness check, we instead use cohorts $i + 6$ to $i + 10$, namely cohorts that joined 6 to 10 months earlier, as the control group. This alternative specification allows us to assess whether our conclusions depend on the exact set of nearby control cohorts used in the baseline TCI construction.

Third, we address the concern that, under the baseline TCI design, treated and control cohorts may share a common calendar-time period during which both are active. Although such overlap does not contaminate the baseline design mechanically, it raises the possibility that the activity of one group may affect the other through channels not fully absorbed by tenure fixed effects. To address this concern, we construct an alternative control group that is further removed in entry time: for treated cohort i , we use cohorts $i + 11$ to $i + 15$, namely cohorts that joined 11 to 15 months earlier. This alternative comparison removes the common period between treated and control cohorts.

Last but not least, we examine whether the comparability of cohorts weakens when treated cohorts are matched to much earlier cohorts that may have faced different expectations, norms, or competitive conditions. To reduce exposure to such long-run calendar-time heterogeneity, we exclude early cohorts and restrict the analysis to the 24 cohorts from the most recent two years that pass the Granger and KS tests.

Taken together, these additional exercises strengthen the plausibility of the identifying assumptions underlying our TCI design and show that the main findings are robust to alternative sample restrictions and control-group constructions.

Table G.7 Sample-Selection Robustness Checks under Temporal Causal Inference Using Difference-in-Differences

Dependent Variable	Lofter				Graffiti Kingdom			
	(1) Exclude Superusers	(2) Control Cohorts ($i + 6$ to $i + 10$)	(3) Control Cohorts ($i + 11$ to $i + 15$)	(4) Exclude Early Cohorts log(Num of Works + 1)	(5) Exclude Superusers	(6) Control Cohorts ($i + 6$ to $i + 10$)	(7) Control Cohorts ($i + 11$ to $i + 15$)	(8) Exclude Early Cohorts
Treat $\times$ Post Month 1	-0.0841*** [0.0073]	-0.0843*** [0.0146]	-0.0906*** [0.0234]	-0.0717*** [0.0090]	0.0239*** [0.0047]	0.0251*** [0.0098]	0.0320*** [0.0114]	0.0295*** [0.0071]
Treat $\times$ Post Month 2	-0.1622*** [0.0071]	-0.1613*** [0.0153]	-0.1538*** [0.0261]	-0.1283*** [0.0087]	0.0221*** [0.0043]	0.0221*** [0.0094]	0.0271** [0.0111]	0.0286*** [0.0063]
Treat $\times$ Post Month 3	-0.1617*** [0.0071]	-0.1552*** [0.0174]	-0.1462*** [0.0270]	-0.1243*** [0.0086]	0.0119*** [0.0042]	0.0099 [0.0079]	0.0125 [0.0101]	0.0184*** [0.0055]
Estimation Method	Difference-in-Differences							
Creator Fixed Effects	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Tenure Fixed Effects	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Number of Creators	138,825	112,429	101,215	69,697	23,985	17,156	13,264	15,127

Note. The unit of analysis is the creator-month. Cohort i denotes the treated cohort in the control-window labels. The “Exclude Superusers” specifications drop the top 1% of creators in each cohort by total work volume, and the “Exclude Early Cohorts” specifications keep the 24 most recent validated treated cohorts. Robust standard errors clustered at the creator level are in brackets. * $p < 0.1$, ** $p < 0.05$, *** $p < 0.01$.

Online Appendix H Text Analysis: Methodology and Validation

General Text Analysis Procedure

The textual analysis of Lofter posts consists of three main steps, summarized as follows. First, we collect all posts from Lofter and then use an LLM (gpt-4o-mini) to (i) distinguish posts that express attitudes or opinions about Generative AI artworks from those that are irrelevant, and (ii) classify the attitude expressed in those posts that do discuss Generative AI. We implement this step using the following prompts:

```
SYSTEM_PROMPT = """\
You are a precise classifier for AI-generated content discussions.
Goal: Decide whether a text expresses an *opinion/judgment* about *AI
creation* (especially AI art/AI drawing), and determine stance.

Definitions:
- OPINION_ON_AI_CREATION: The text conveys the author's opinion or
  evaluation (support/oppose/mixed/neutral) about AI creation. It may
  discuss ethics, copyright, plagiarism, creativity, skill, job market
  , fairness in contests, model/regulation, value judgment, personal
  attitude.
- NON_OPINION_AI: AI-creation-related but no opinion: tutorials,
  factual Q&A, news relay, tool setup, pure description without stance
  .
- IRRELEVANT_OR_FICTION: Unrelated to AI creation; or clear fiction/
  novel/fanfic/roleplay/character narratives, even if "AI" appears
  incidentally.

Output JSON fields (STRICT):
- label: one of ["OPINION_ON_AI_CREATION", "NON_OPINION_AI", "
  IRRELEVANT_OR_FICTION"]
- stance: one of ["pro", "anti", "neutral", "mixed", "n/a"] # 'n/a' if not
  OPINION_ON_AI_CREATION
- rationale_zh: <= 40 chars, brief reason
- opinion_snippet: a short quote (<= 40 chars) that best evidences the
  judgment; empty if none

Be conservative: if mainly fiction/novel, choose IRRELEVANT_OR_FICTION.
If AI-related but purely instructional/news without stance, choose
NON_OPINION_AI.
"""
```

Based on the outputs, we can classify the relevant posts into four categories: (i) dislike AI; (ii) pro AI; (iii) neutral; and (iv) mixed feelings. Then we use the LLM again to extract

the detailed reasons why the content creators dislike Generative AI, with the following prompts:

```
SYSTEM_PROMPT = """\
You are a precise annotator. The input texts come from people who
    DISLIKE or OPPOSE AI creation (especially AI art/drawing).
Task: Extract the KEY REASONS (free-form, no fixed taxonomy).

Output STRICT JSON:
- reasons: array of 1-5 concise English phrases (2-8 words each)
    capturing the main reasons (e.g., copyright infringement, job
    displacement, unfair contests, low quality, privacy/portrait rights,
    opaque training, deepfake risk, censorship, high learning cost).
    Avoid redundancy.
- evidence_snippet: <= 40 characters quoting the best evidence from the
    text.
- confidence: number in [0,1].

Guidelines:
- Include only reasons for DISLIKING/OPPOSING AI creation. Be specific
    but concise.
- Prefer stable, generalizable phrases (avoid overly specific names/
    events).
- If no clear reason, return [] with low confidence.
"""
```

Last but not least, for each of the top three broad reasons (Replacement Risk, Low Quality, and Copyright Infringement), we apply the following prompt to extract more detailed and concrete underlying reasons:

```
SYSTEM_PROMPT = """\
You are a rigorous annotator. Focusing only on the given broad category
    {category}, please extract no more than six specific reasons (short
    Chinese phrases) from the following posts, and provide the
    occurrence count (an integer) for each reason.

Return only a strict JSON object.
"""
```

Categorizing Posts by Target (March 2023 Sub-Analysis)

For the separate analysis discussed in Section 4.3, we restrict the sample to posts expressing negative attitudes toward Generative AI during the period in March 2023 slightly before

and after the Generative AI feature was removed and classify them into three categories by what they primarily target:

- **(A) Directly feature-related:** posts that explicitly mention the “Profile Picture Generator” or refer to it through indirect phrases like “this tool” when the surrounding context clearly refers to the Generative AI feature. We rely on keyword searches for explicit references such as “profile picture generator” and “drawing machine”, together with transparent keyword and context rules for indirect references such as “this tool” when the surrounding context clearly refers to the Generative AI feature.
- **(B) Platform-related but not feature-specific:** posts that mention the platform (using keywords such as “Lofter,” or implicit expressions such as “the platform” when the discussion clearly concerns Lofter’s deployment of the AI feature) but do not mention the feature itself.
- **(C) General negative attitudes:** posts that express general negative attitudes toward Generative AI without referring to either the feature or the platform.

When a post mentions both the tool and the platform in the same discussion, we classify it as category (A) rather than (B). Any post that does not meet the criteria for (A) or (B) is classified as (C).

Validation against Human Annotators

To verify the accuracy of our LLM-based classification, we recruited eleven human annotators to code a sample of posts drawn from NetEase Lofter. All annotators are native Chinese speakers and graduate students enrolled in top art academies in China, and they are active users of Lofter who are familiar with the platform’s communities, norms, and terminology. The three annotation tasks are designed to validate our LLM-based classifications. The annotation tasks were organized into three datasets (each containing approximately 100 posts), and annotators were instructed to assign exactly one label per post unless otherwise specified. The coding instructions for each dataset are summarized below.

Dataset 1: Whether the post expresses an attitude toward AI art

Task. Determine whether the post explicitly expresses an attitude toward “AI Drawing/AI Art”.

Labels.

- **1 = Yes:** The post clearly expresses a view or attitude toward AI Drawing/AI Art (e.g., liking, supporting, opposing, or worrying about AI Drawing/AI Art).
- **0 = No:** The post may mention “AI” or related content, but does not express an attitude toward AI Drawing/AI Art, or is unrelated to AI Drawing/AI Art.

Dataset 2: Overall stance toward AI Drawing/AI Art

Task. All posts in this dataset are confirmed to be about AI Drawing/AI Art. Judge the overall affective stance toward AI Drawing/AI Art.

Labels.

- **1 = Positive:** Clearly expresses liking, support, approval, anticipation, or other positive attitudes.
- **2 = Negative:** Clearly expresses dislike, opposition, concern, dissatisfaction, or other negative attitudes.
- **3 = Mixed:** Contains both clearly positive and clearly negative evaluations so that it cannot be simply classified as positive or negative.
- **4 = Neutral:** Describes facts, asks questions, or discusses issues in an objective tone without a clear affective stance.

Dataset 3: Specific reasons for disliking AI Drawing/AI Art

Task. All posts in this dataset express dislike of or opposition to AI Drawing/AI Art. Identify the primary reason for disliking or opposing AI Drawing/AI Art in the post.

Labels.

- **1 = Replacement Risk:** Concern that AI Drawing/AI Art will replace human artists, affecting jobs, income, or survival space.
- **2 = Low quality:** Belief that AI works are of poor quality (e.g., stiff style, anatomical errors, lack of emotion).

- **3 = Copyright infringement:** Concern or accusations that AI training or generation involves infringement (e.g., unauthorized use of others' works, "art theft", misattribution, or difficulty protecting original rights).

- **4 = Other:** Reasons that do not fall into the above categories or are too vague/ambiguous to classify.

After we collect the annotated data, we first assess agreement across the eleven human annotators and, for each dataset, derive a single gold label for every post using a simple majority-voting rule. We then compare the LLM's predictions to these human gold labels.

For Datasets 1–3, Fleiss' κ for the human annotators is 0.77, 0.55, and 0.68, respectively, indicating moderate to substantial agreement among coders on these subjective classification tasks. Relative to the resulting gold labels (Fleiss 1971, Landis and Koch 1977), the LLM achieves accuracies of 0.94, 0.92, and 0.81 in Datasets 1–3, respectively. Thus, for the first two tasks the LLM's classifications are very close to the consensus human coding, and even for the more fine-grained reason taxonomy in Dataset 3, the LLM attains a level of agreement with the gold standard that is comparable to typical human–human reliability in similar multi-category content-analysis settings.

References

- Babina T, Fedyk A, He A, Hodson J (2024) Artificial intelligence, firm growth, and product innovation. Journal of Financial Economics 151:103745.
- Burtch G, Lee D, Chen Z (2023) The consequences of generative ai for ugc and online community engagement. Available at SSRN 4521754 .
- Cavusoglu H, Phan TQ, Cavusoglu H, Airoidi EM (2016) Assessing the impact of granular privacy controls on content sharing and disclosure on facebook. Information Systems Research 27(4):848–879.
- Chen J, Roth J (2024) Logs with zeros? some problems and solutions. The Quarterly Journal of Economics 139(2):891–936.
- Dixon WJ, Massey Jr FJ (1951) Introduction to statistical analysis. .
- Eichenbaum M, Godinho de Matos M, Lima F, Rebelo S, Trabandt M (2024) Expectations, infections, and economic activity. Journal of Political Economy 132(8):2571–2611.
- Fleiss JL (1971) Measuring nominal scale agreement among many raters. Psychological bulletin 76(5):378.
- Goldberg SG, Johnson GA, Shriver SK (2024) Regulating privacy online: An economic evaluation of the gdpr. American Economic Journal: Economic Policy 16(1):325–358.
- Gourieroux C, Monfort A, Trognon A (1984a) Pseudo maximum likelihood methods: Applications to poisson models. Econometrica: Journal of the Econometric Society 701–720.
- Gourieroux C, Monfort A, Trognon A (1984b) Pseudo maximum likelihood methods: Theory. Econometrica: journal of the Econometric Society 681–700.
- Granger CW (1980) Testing for causality: A personal viewpoint. Journal of Economic Dynamics and control 2:329–352.
- Hausman C, Rapson DS (2018) Regression discontinuity in time: Considerations for empirical applications. Annual Review of Resource Economics 10:533–552.
- Landis JR, Koch GG (1977) The measurement of observer agreement for categorical data. biometrics 159–174.
- Rubin DB (1980) Randomization analysis of experimental data: The fisher randomization test comment. Journal of the American statistical association 75(371):591–593.
- Shi Z, Liu X, Srinivasan K (2022) Hype news diffusion and risk of misinformation: the oz effect in health care. Journal of Marketing Research 59(2):327–352.

Silva JS, Tenreiro S (2006) The log of gravity. The Review of Economics and statistics 641–658.

Turjeman D, Feinberg FM (2023) When the data are out: measuring behavioral changes following a data breach. Marketing Science .

Xu Y (2017) Generalized synthetic control method: Causal inference with interactive fixed effects models. Political Analysis 25(1):57–76.