

Episode 1 Blueprint: The Conversation

Audit Prompt Chain

Source workflow demonstrated by Brooke Sellas, B Squared Media, on The AI Social Playbook episode 1. Every prompt below is reproduced from the recording. Reconstructed text is marked where the on-screen prompt was not read aloud in full.

Technical stack

Tool	Role	Configuration observed
Claude	Analysis engine for all four passes	Standard web chat. No project, no custom skill. Model version and plan tier were not stated on air.
Social media management tool or native platform analytics	Source of the conversation export	Agorapulse named as the example. Native channel analytics stated as the equivalent path. Report lives under inbox or messages depending on platform.
Spreadsheet or CSV	Transport format and verification index	One month of data per file. Contains live hyperlinks to each source conversation.
Prompt library	Storage for the reusable prompts	Location and tool not specified. Any store the operator can copy from will

		do.
--	--	-----

ChatGPT, Gemini, and Copilot were named in the episode only as platform equivalents for storing standing instructions. They are not part of this workflow.

Configuration

Chat, not project. The demonstration ran in a plain chat rather than a saved project. The stated reason was to keep stored assumptions from shaping a live analysis. An operator who prefers a project should expect the project's instructions to influence theme selection and should account for that.

Account history affects output. Spam was removed from the dataset without any instruction to do so, attributed to repeated prior use of the same workflow on the same account. A first-time operator will not get this behaviour and must request spam removal explicitly in the first prompt.

Export file structure. One calendar month per file. Required fields observed in the demonstration: message text, source platform, surface type (organic comment or ad comment), and a hyperlink to the live conversation. The hyperlink field is what makes stage 4 verification fast; without it, verification still runs but requires manual search.

Data volume. The demonstration dataset held 1,902 incoming messages before spam removal and 1,829 after. Processing time was reported as a few minutes per pass at that volume.

Workflow map

Stage	Input	Action	Output
1. Export	One month of inbox activity	Pull the conversation report from the social tool or native analytics	CSV with message text, platform, surface type, permalink
2. Theme pass	CSV plus Prompt 1	Categorize by theme with frequency, tone,	Ranked themes with mention counts, percentage share,

		and stated intent	and positive/negative split
3. Complaint pass	Stage 2 output plus Prompt 2	Rank negative themes by volume and by one-sidedness	Top three complaints with negative mention counts and share of total negative volume
4. Verification pass	Stage 3 output plus Prompt 3	Retrieve the verbatim comments behind each ranked complaint with their file locations	Quoted messages with row references, checked by a human against the source
5. FAQ pass	Same conversation plus Prompt 4	Extract repeated questions and match each to a content format	Ranked question list with volume and a suggested zero-click format per question
6. Reporting	Verified stage 3 and stage 5 output	Collapse into good and bad, attach proposed fixes to each friction point	Client-facing summary and a content brief

Prompt 1: context and theme categorization

Read aloud from the screen. Reproduced as spoken, including the disfluency in the final sentence.

You're analyzing customer conversations from a kitty litter brand's social media channels. This brand sells kitty litter to discerning cat owners.

The following data includes comments, DMs, replies from July 2026. Your job is to help me understand what customers are actually saying, feeling, and needing.

Categorize these conversations by theme. For each theme, tell me how frequently appears the emotional tone and what the customer is actually asking for or trying to accomplish.

The final sentence as delivered compresses three requests into one clause. The intended asks, stated elsewhere in the episode, are frequency of the theme, the emotional tone of the theme, and what the customer wants. A cleaned version follows. It is a reconstruction, not a transcript quote.

You're analyzing customer conversations from a {BRAND_CATEGORY} brand's social media channels. This brand sells {PRODUCT_DESCRIPTION} to {BUYER_DESCRIPTOR}.

The following data includes {DATA_TYPES} from {DATE_RANGE}. Strip out spam before you begin and tell me how many records you removed. Your job is to help me understand what customers are actually saying, feeling, and needing.

Categorize these conversations by theme. For each theme, tell me how often it appears as a count and as a percentage of total conversations, the emotional tone split, and what the customer is actually asking for or trying to accomplish.

Prompt 2: complaint ranking

Described on air rather than read from the screen. Reconstructed from the spoken description and from the model's own restatement of the task.

What are our top three negative complaints worth addressing? Rank them by raw negative volume and by how one-sided the complaint is, not just by theme size. For each one give me the negative mention count and its share of total negative volume for the period.

The ranking logic is the reason this pass exists. In the demonstration it returned price and value objections at 100 negative mentions and 17.2% of negative volume, dust at 66 negative mentions where 55% of all dust mentions were negative, and clumping at 53 negative mentions flagged as the most reputationally damaging. Dust outranked larger themes because of its lopsidedness, not its size.

Prompt 3: verification

Spoken close to verbatim during the demonstration.

Bring up the actual comments that you're flagging here alongside my top three complaints.

Add row and column references when the file is large. The operator can also fold this requirement into Prompt 1 so the evidence arrives with the first analysis. That variant, proposed by the host during the episode, runs:

Categorize this and give me three examples of actual comments from each category, and tell me what row and column they are in.

Expect the model to offer to draft customer replies at this stage. Decline and continue. The workflow is an analysis pass, and reply drafting pulls the session off task.

Prompt 4: FAQ extraction and format matching

Described on air rather than read from the screen. Reconstructed from the spoken description and the model's restatement, which was identifying trending FAQs for social media content creation.

What are the top questions people asked this month? Pull them from actual repeated questions in the data, rank them by how often each was asked, and give me the volume for each. Then give me a content-ready summary that pairs every question with the best zero-click format for answering it.

Output in the demonstration was a ranked question list covering litter box compatibility, cost, subscription mechanics, clumping, tracking outside the box, competitor comparison, and odor control. The content-ready summary matched formats to question types: a compatibility graphic, a cost breakdown carousel, and a short demo video or GIF. The competitor comparison question surfaced named rival brands from inside the customer comments, which is the input for a comparison landing page.

Variables to replace

Variable	What it is	Worked example
{BRAND_CATEGORY}	The product category, not the company name	kitty litter
{PRODUCT_DESCRIPTION}	What the brand sells, with the differentiator that matters to buyers	sustainable kitty litter
{BUYER_DESCRIPTOR}	The buyer in the brand's own words, not a demographic bracket	discerning cat owners
{DATA_TYPES}	Every surface represented in the export	comments, DMs, replies, ad comments
{DATE_RANGE}	The single period the file covers, stated explicitly so	July 2026

	the model does not infer it	
--	-----------------------------	--

Human quality control, three checks before anything ships

1. Reconcile the record counts across every pass. The theme pass reports a pre-spam total and a post-spam total. Later passes may cite either one. Confirm which number each conclusion rests on before a percentage goes into a client deck. In the source episode the two figures were 1,902 and 1,829, and both appeared in reference to the same analysis. This check catches percentage errors that survive every other review because the arithmetic is internally consistent against the wrong denominator.

2. Open the source rows behind every ranked complaint. Read at least three verbatim comments per ranked theme against the live conversation, using the permalink column. This check catches sarcasm scored as sincere complaint and clusters the model assembled from loosely related messages. It is the only step that tests whether the ranking describes the data or describes a pattern the model constructed.

3. Confirm the FAQ list contains repeats, not one-offs. Ask for the count behind each question and reject any question that appeared fewer than a handful of times. A low-volume dataset will return a list of distinct questions asked once each, which looks identical to a real FAQ ranking and is worthless as a content calendar. This check catches a false content brief before a month of production goes into it.