

Prompting Techniques

Technique

AI Focus Group

A way to pressure-test a finished or near-finished ad against the people you're hoping will respond to it, before you spend on placement. The prompt runs three or four named personas through your ad, surfaces where their reactions agree and diverge, and ends with a one-line verdict from each: sign up, hesitate, or pass, and why. Make the personas disagree, give the model only what a consumer would know about the brand (five lines, not the campaign brief), and treat the output as hypotheses to test, not fixes to apply. Run it after the ad is built but before you publish, alongside a brief scoring prompt that critiques the same ad from the inside. The two together tell you whether the ad lands as communication and delivers what the brief said it should.

Transcript Generator

The Transcript Generator turns your finished ad into a shot-by-shot table, capturing every visual, voiceover line and on-screen text with precise timing. Its role is to give the brief scorer something readable to work with. If you're running the brief scorer in a text/image-only model such as Claude, this transcript is a prerequisite: it is what the scorer holds up against your Guidance Summary to check for hook drift, missing pillars or CTA fidelity. If you're running the brief scorer in Gemini, you can skip this step and feed the video directly, since Gemini can analyse video natively. Same output either way, one fewer step

Example

Act as a marketing research focus group reviewing a short-form **[PLATFORM, e.g. Instagram]** video ad for **[BRAND NAME]**.

[BRAND NAME] in the world (what a consumer would know)

[ONE TO FIVE LINES OF CONSUMER-GRADE BRAND DESCRIPTION. Cover: what the brand is, what the first product or service is, what the brand stands for in plain language, the stage the brand is at (pre-launch, recently launched, established), and what the ad's call to action is asking the viewer to do. This is consumer-grade context only, not the campaign brief, not the messaging architecture, not internal documents. The personas should know only what a consumer would learn from the ad and these few lines.]

Inputs attached

1. The video ad.
2. A focus group personas document. Use these personas exactly as written. Do not invent new participants, do not soften or merge personas, and keep their distinct voices throughout.

Task

Run a realistic back-and-forth group discussion in which each persona reacts to the ad. They should challenge one another, surface doubts, build on each other's points, and disagree where their priorities differ. Keep the conversation natural and detailed, with clearly distinct voices.

Cover, in roughly this order:

1. First impressions of the opening three seconds: opening shot, voiceover hook, on-screen text. Where did each persona's attention go, and did anything make them want to scroll past.
2. Clarity of the core message. What do they each think the ad is asking them to believe and to do.
3. Credibility and trust signals. Which claims feel earned, which feel thin, and what specific proof is missing for them to take the brand seriously at this stage.
4. Emotional response and brand fit. Does the tone feel right for the kind of **[CATEGORY, e.g. skincare / fitness / fintech]** brand each persona buys from. What does the ad remind them of, positively or negatively.
5. Reaction to the call to action specifically. Is the CTA clear? The ad asks them to **[DESCRIBE THE ACTION THE CTA IS ASKING FOR, e.g. join a waitlist, book a demo, start a free trial, buy the product]**. Would each persona **[TAKE THE ACTION]**, hesitate, or pass, and why. Distinguish "I would tap through to read the page" from "I would actually **[COMMIT, e.g. submit my email, book the call, enter my card details]**".
6. Specific improvements. Each persona names one change to the ad that would move them closer to **[THE DESIRED ACTION]**. Be concrete: a line of voiceover, a shot, a specific piece of proof, a change to the CTA.

Rules

- Personas only know what a consumer would know from watching the ad and the short brand description above. They do not know the campaign brief, the messaging architecture, the success benchmarks or any internal brand documents. Do not have them reference language they could not plausibly have access to.
- Do not invent claims, certifications, ingredients, partnerships or reviews that are not visible or audible in the ad.
- Keep at least one persona open to **[THE DESIRED ACTION]**. The group should surface tension, not collapse into uniform scepticism.
- Use **[British / American / Canadian]** English.
- End with a short summary section labelled "Headline verdicts", one line per persona, in the form: "[Name]: [sign up / hesitate / pass] because [one-sentence reason]."

Note: Remember to attach your campaign video and your personas document.

Act as a video transcription assistant for short-form advertising.

I'm attaching a short-form ad video (under 30 seconds). Produce a shot-by-shot transcript as a table with these columns:

· Beat · Timing · Visual · On-screen copy · Voiceover

Rules:

- One row per shot. A new shot is any visible cut or scene change.
- Beat: label each shot with its narrative function (Hook, Problem, Proof, Payoff, CTA, or similar).
- Timing: start–end in seconds to one decimal place, e.g. 5.0–7.0s.
- Visual: one sentence describing what is on screen. Be concrete: what is the object, who is in shot, what is the lighting.
- On-screen copy: any text shown on the frame, verbatim. If none, "—".
- Voiceover: the spoken line during that shot, verbatim. If none, "[ambient beat]" or "[silent beat]"

Do not invent on-screen copy or voiceover that is not actually present in the video. If you cannot make out a line, write "[unclear]" rather than guessing.

At the end, state the total runtime in seconds.

Prompting Techniques

Technique

Brief Scorer

A way to pressure-test a finished ad against the brief that produced it, to find where the brief itself needs to sharpen for the next cycle. The prompt scores the ad across nine dimensions (audience, single message, hook, proof, CTA, conversion readiness, tone, visual direction, banned content), each as green, amber or red, with a "fix in brief" or "fix in ad" note where the score isn't green. Run it with your latest Campaign Guidance Summary and either the ad video (vision-capable models) or a shot-by-shot transcript. The goal is to surface upstream gaps, most fixes should resolve to the brief rather than the ad, because the brief is the system everything downstream keys off. Pair it with the AI focus group prompt: the scorer critiques the ad from the inside, the focus group critiques it from the outside, and the two together tell you whether the ad lands as communication and delivers what the brief said it should.

Example

Role

Act as an expert campaign critic and evaluator. You score short-form video ads against the brief that produced them, with the goal of surfacing where the brief itself needs to sharpen for the next cycle. Score the ad against the brief, not against your own opinion of what makes a good ad.

Inputs

You have two attached documents:

1. The Campaign Guidance Summary (V2). This contains the campaign positioning, audience, messaging, tone of voice, visual direction, conversion readiness criteria and quality checks for the brand.
2. The Ad Transcript. This is the audience-facing record of the finished ad, as a shot-by-shot table with columns for shot number, beat, timing, visual, on-screen copy and voiceover.

The transcript is the canonical record of what the audience actually experiences. Do not work from any other artefact (storyboard, brief, voiceover-only text). If the transcript is missing any column, flag that as a constraint before scoring.

Your job

Score the ad against the brief, identify where each delivers, falls short, or fails. The goal is not to score the ad in isolation. The goal is to surface where the brief itself needs to sharpen for the next campaign cycle, so the iteration loop can feed forward.

Score the ad across these nine dimensions:

1. Audience match. Does the ad land for the primary audience defined in the brief?
2. Single message. Is the one core message defined in the brief delivered clearly within the first 5 seconds?
3. Hook strength. Does the opening shot and voiceover hook stop the scroll for the defined audience?
4. Proof. Are claims specific, verifiable and supported by what is visible or audible in the ad?
5. CTA alignment. Does the CTA match the brief's stated CTA ceiling and feel earned by what came before?
6. Conversion readiness. Does the CTA point to a real conversion asset, with a defined conversion event, as declared in the brief?
7. Tone of voice. Does the ad sound like the brand's defined voice, including correct use of the words-to-use and words-to-avoid lists?
8. Visual direction. Does the visual treatment match the brief's defined mood, palette, composition guidance and what-to-avoid notes?
9. Banned content and substantiation. Does the ad use any words, phrases or claims the brief explicitly tells us to avoid, or make any claim that would need substantiation under UK advertising rules?

For each dimension

- Score as green (delivers), amber (partial or unclear), or red (fails or absent).
- Give a one-line verdict explaining the score, referencing specific shots by number or brief sections by name.
- Where amber or red, write a one-line "fix in brief" note that names the specific brief line or section that needs to change in the next iteration. If the gap is in the ad rather than the brief, write "fix in ad" instead. Most gaps should resolve to "fix in brief", the brief is the upstream system.

Rules

- Quote directly from the brief and from the transcript when justifying a score. Use the format: brief says "...", transcript shows "..." (beat n).
- Do not invent details that are not in the attached documents.
- Do not soften scores to be diplomatic. An amber that should be red costs the next campaign cycle.
- Use **[British / American / Canadian English]**.

Consistency check

Before generating the HTML, run a consistency check. For each dimension, confirm the status matches the verdict and the verdict references specific evidence from the attached documents. Correct any mismatch before continuing.

Output

Output the final scorecard as a single HTML artefact styled as a one-page brief scoring dashboard. Layout reads top to bottom like a one-page report.

Title the dashboard:

"[BRAND NAME] · [AD VERSION] · Brief Scoring Dashboard".

- Header: campaign name, ad version, date scored, overall score (count of green / amber / red across the nine dimensions).
- Body: one card per dimension, displaying the dimension name, a status indicator (green / amber / red dot or pill), the one-line verdict, and the fix-in-brief or fix-in-ad note where present.
- Footer: a "Carry into next brief" panel listing the specific brief lines or sections to rewrite, drawn from the fix-in-brief notes above. Group by brief section if more than three.

The HTML must be a direct rendering of the final scorecard. Each card must display the same status, verdict and fix note assigned in the scoring pass. Do not reinterpret or reassign at the HTML stage.

Prompting Techniques

Technique

The Performance Analyst

A way to read campaign performance data against the brief that produced it, instead of in isolation. The prompt scores six dimensions (cost efficiency, hook, landing page, placement mix, day patterns, creative health) and splits findings into short-loop fixes for this campaign now and long-loop fixes for the next. Pair it with the brief scorer and AI focus group, all three plug into the same V3 Guidance Summary and produce findings that feed forward into the next brief.

Example

Role

Act as a senior paid social performance analyst. You score campaign performance against the brief that produced it, with the goal of identifying both the immediate levers for this campaign and the structural improvements for the next.

Inputs

You have three attached documents:

1. The campaign performance data (Ads Manager export, screenshot, or table covering at least cost, impressions, hook rate, clicks, conversions and placement breakdown).
2. The Campaign Guidance Summary, containing the conversion readiness criteria and benchmarks the campaign was held to.
3. (Optional) The Ad Transcript or finished video, so performance signals can be connected to specific shots or beats.
- 4 (Optional): A landing page mock up prompting the target audience to e.g. sign up, buy, etc.

Your job

Read the data against the brief's benchmarks, not in isolation. Identify where each benchmark hit, missed, or showed a near-miss, and explain why using evidence from the data and the ad.

Surface levers for two horizons:

- Short loop: improvements that close gaps in THIS campaign now.
- Long loop: improvements to the brief and prompt stack that would sharpen the NEXT campaign for this brand.

Analyse the data across these five dimensions:

1. Cost efficiency. Is cost per conversion within the brief's benchmark band? Where it isn't, name the leak with specific numbers.
2. Hook performance. Does the hook rate clear the brief's target? Read the daily pattern to surface fatigue.
3. Placement mix. Which placements drove which volume at what cost. Identify outliers and underperformers.
4. Day patterns and fatigue. How metrics evolved D1 to D7. Look for the learning phase, peak, and decay.
5. Creative health. Is the ad earning attention at a cost that matches what each placement charges?

For each dimension

- Score as green (delivers), amber (partial or near-miss), or red (fails or absent).
- Give a one-line verdict with specific numbers quoted from the data.
- Distinguish short-loop fixes from long-loop fixes.

Rules

- Quote specific numbers from the data and specific lines from the brief when justifying findings.
- Distinguish causation from correlation. If two metrics moved together, say so. Do not assume one caused the other without evidence.
- Do not invent metrics or data that are not in the attached documents.
- Use **[British / American / Canadian]** English.

Output

Output the analysis as a structured report with the following sections:

1. Headline. One sentence summarising the campaign verdict.
2. Benchmark scorecard. Each of the six dimensions with its green/amber/red score and one-line verdict.
3. The leak. The single biggest cost driver, named explicitly with the number that proves it.
4. Improve this campaign now (short loop). Three to five concrete actions, each with the specific change and the metric it would move.
5. Improve future campaigns for this brand (long loop). Three to five brief or prompt-stack improvements, named against specific brief sections.
6. The top lever. The single change that would move the needle most. State it as: "Apply [X], expect [Y]."

Prompting Techniques

Bonus Prompt

Technique

CTA Audit

Run your CTA line through a structured four-question test against your Campaign Guidance Summary to check whether it is honest, achievable and aligned to your brand's actual conversion readiness. The auditor scores audience fit, brand stage, tone commitment and realistic next action as pass or fail, then flags any brief-level gaps the Guidance Summary failed to declare. Use it as a fast diagnostic before you finalise a brief, or as a post-hoc check when a CTA feels off but you cannot yet name why. Outputs a verdict headline, a four-question scorecard and a specific fix-in-brief note pointing to the section or prompt that needs to change. Tool: ChatGPT or Claude. Inputs: your CTA line plus your Campaign Guidance Summary.

Example

Act as an expert campaign brief auditor.

You are auditing the CTA in a short-form video ad against the brief that produced it. Your job is not to critique the ad. Your job is to test whether the brief gave the CTA everything it needed to be honest, achievable and aligned to the brand's actual conversion readiness.

Inputs attached

1. The Campaign Guidance Summary.
2. The current CTA line from the video ad.

Task

Run the CTA through a four-question test.

1. Audience fit. Who is the audience, and how do they buy in this category? Does the CTA match how they actually move from awareness to purchase?
2. Brand stage. What stage is the brand at? Does the CTA point to an action the brand can currently support with a live asset?
3. Tone commitment. What tone does the brief commit to? Does the CTA hold that tone, or does it drift?
4. Realistic next action. What is the realistic next action this ad can drive for a first-time viewer, given brand stage and audience buying behaviour? Does the CTA match that action, or does it overreach?

Score each question as pass or fail with a one-line verdict, quoting from the Guidance Summary where relevant.

Then answer three brief-level questions: -

- Does the brief declare what conversion assets the brand has live?
- Does the brief state the strongest honest CTA the brand can support?
- Does the brief define what "this ad worked" actually means?

Rules

- Do not soften a fail to be diplomatic. A soft fail costs the next campaign cycle.
- Quote directly from the Guidance Summary when justifying a verdict. Use the format: brief says "...".
- Do not invent brand assets, audiences or benchmarks not present in the attached document.
- Use British / American / Canadian English.

Output

1. Verdict headline. One sentence: does the CTA pass or fail, and where does the failure sit, in the CTA or in the brief.
2. Four-question scorecard. Pass or fail, with one-line verdict for each.
3. Brief gaps. Any of the three brief-level questions the Guidance Summary does not answer.
4. Fix in brief. Name the specific brief section or prompt that needs to change so the CTA cannot invent itself again.