



## The Rundown

# Build a Marketing Creative Studio in your Pocket

By Melanie Moeller

Founder and Chief AI Officer, GenFutures Lab



# Session 2 recap

---

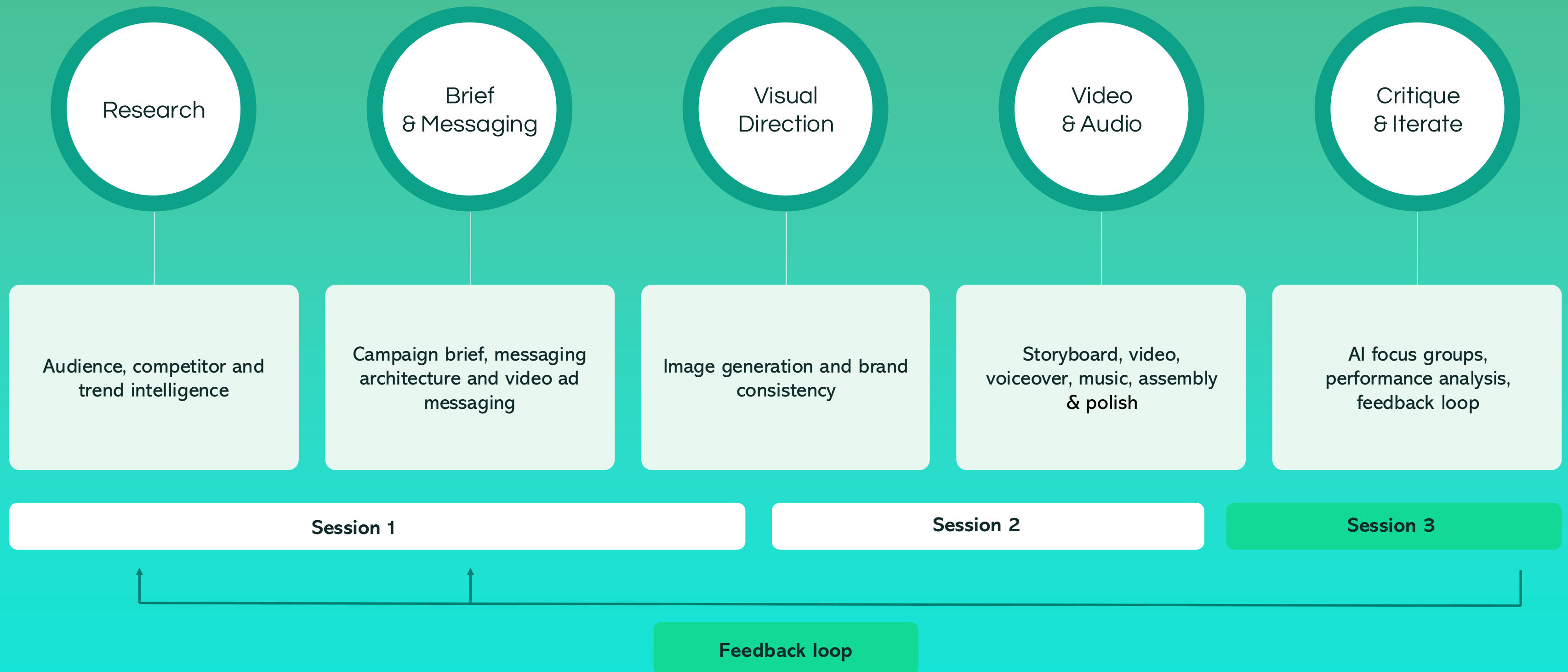
## What you built in Session 2

- A shot-by-shot storyboard mapped to the five-beat narrative arc.
- Short video clips animated from your campaign images.
- A voiceover and background music tuned to your brand and pacing.
- A rough cut of your campaign video, assembled in CapCut.

## What to have with you today

- **Your Campaign Guidance Summary.** The portable brief that feeds every prompt today.
- **Your V1 rough cut.** The ad we'll reversion, critique, test and improve in this session.

# AI workflow



*The feedback loop runs from Critique and Iterate back to Research and Brief. That loop is the whole of Session 3.*

# Session 3

Test, Polish and Build the Feedback Loop

TODAY, IN THREE LOOPS

# How today works

Each loop tests the ad against something different.

01

## CTA optimisation

Surface the gap. Update the prompts that build the brief.



02

## AI critique

Stress-test the V2 ad before we simulate running the ad.

---

2a · Personas  
2b · Brief scoring



03

## Performance analysis

Look at the data. Decide the creative moves. Feed both back.

***The principle:** We don't just run each loop. We feed what it surfaces back in before running the next. That is what makes today an engine, not just a test.*

# CTA Optimisation

## Loop 1



# What is your gut reaction to these CTA's?

Where does that reaction come from?

BRIEF-BASED CTA

“Read the choices behind Sustain Well.”

MY CTA · PATCHED IN THE SESSION

“Join the Sustain Well circle.”

# The patch fixed the feeling, not the brief

---



Does Sustain Well have a circle to join?

No. No community asset exists at concept stage.



Does the audience convert on community joins?

No. They read before they buy.



Was the patch tested against the brief?

No. It was an instinct fix.

---

*The discomfort was a real signal. But the cause was upstream, not in the script.*

# The CTA passed. The brief didn't.

## FOUR-QUESTION TEST, AND WHAT IT MISSED

### Four-question test

- 1 Who is the audience, and how do they buy in this category?
- 2 What stage is the brand at?
- 3 What tone does the brief commit to?
- 4 What is the realistic next action this ad can drive?

*"Read the choices behind Sustain Well"*

Passes all four

*"Join the Sustain Well circle"*

Fails three of four

### And here's what the brief didn't asked

- X What conversion assets the brand has live
- X The strongest honest CTA the brand can support
- X What "this ad worked" actually means

To fix the CTA permanently, we update the prompts that build the brief.

# The V2 prompt stack, in four updates

## PROMPT 1 OF 4

### Research: Campaign Intelligence Engine

*ChatGPT Deep Research*

NEW

Scope: research and analyse



Conversion behaviour and brand readiness. How the audience moves from awareness to purchase and any gap between expectations and readiness.

## PROMPT 2 OF 4

### Brief & Creative Direction

*ChatGPT or Claude*

NEW

Part 1: campaign brief



Conversion readiness. List conversion assets live, state the CTA ceiling, define one conversion event, set a target benchmark.

LOAD-BEARING

## PROMPT 3 OF 4

### Video Ad Messaging

*ChatGPT or Claude*

MODIFIED

Deliverables 4 and 6



CTA line must match the declared CTA ceiling. Quality check confirms the CTA points to an asset that currently exists.

## PROMPT 4 OF 4

### Campaign Guidance Summary

*ChatGPT or Claude*

NEW

New section 5



Conversion readiness and success metric: conversion asset the CTA points to, defined conversion event, target benchmark and CTA ceiling.

All four Session 1 prompts gain one small, deliberate change. The pattern is the same: declare what the brand can actually deliver, then stop the CTA inventing itself.

Our V2



LOOP 2

# AI Critique

Target Personas  
Simulation

2a

Brief Scoring

2b

# AI focus group · reactions before publishing

THE THREE PERSONAS · UK 25 TO 42

Hannah, 34 · Bristol

The pragmatic sceptic. Buys established sustainable brands, wants external trust signals before she trusts a new one.

Priya, 29 · London

The ingredient researcher. Checks INCI lists and concentrations, treats a waitlist as a watch list until reviews land.

Imogen, 27 · Brighton

The brand-world purist. Screens hard on tone and craft, the most likely to sign up if the brand world holds.

**Inputs** the V2 video and a short consumer-grade brand description. **Not the full Guidance Summary, the personas must not see the brief.** *Keep at least one persona open to signing up, or the group collapses into uniform scepticism.*



# AI Focus Group



Act as a marketing research focus group reviewing a short-form Instagram video ad for Sustain Well.

### Consumer-grade brand description:

*Sustain Well is a pre-launch botanical skincare brand from the UK. The first product is a botanical hydration cream with chamomile and rosemary. The brand emphasises clear ingredient and packaging choices over green marketing claims. No product is on sale yet. The ad's call to action invites viewers to read the brand's ingredient and packaging detail page and join the waitlist to be first to try it when it launches.*

### Inputs attached

1. The V2 ad video.
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 skincare 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 join a waitlist for a brand that has not launched yet. Would each persona sign up, hesitate, or pass, and why. Distinguish "I would tap through to read the page" from "I would actually submit my email".
6. Specific improvements. Each persona names one change to the ad that would move them closer to signing up. 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 signing up. 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 you campaign video and your personas document!**

# Transcription Generator

Tool: Gemini 3.5



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.



# The brief scorer

Claude · inputs: V2 Guidance Summary + ad transcript · output: a one-page HTML dashboard

## SCORE THE AD AGAINST THE BRIEF ACROSS NINE DIMENSIONS

1. Audience match

2. Single message in 5s

3. Hook strength

4. Proof

5. CTA alignment

6. Conversion readiness

7. Tone of voice

8. Visual direction

9. Banned content

Each dimension scores green, amber or red, with a one-line verdict and a fix-in-brief or fix-in-ad note.

# The brief scorer

*NOTE: Make sure to attach your Campaign Guidance Summary and your Ad Transcript before running this prompt.*



**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.

# What each lens saw

## LOOP 2 · MAPPING YOUR OWN FINDINGS

### Focus group only

#### OUTSIDE-IN · WHAT YOUR AUDIENCE FELT

[Persona 1]. [What they flagged in the ad that the brief never asked it to do.]

[Persona 2]. [Same shape. Be concrete: a missing proof, a tonal slip, a claim that did not land.]

### Both agree

STRONGEST  
SIGNAL

#### CROSS-VALIDATED · TWO LENSES

*“[Your shared phrase]”*

**Audience.** [How it felt in the ad, in audience-grade words.]

**Scorer, brief.** [Which brief section or anchor it drifts from.]

### Brief scorer only

#### INSIDE-OUT · BRIEF ALIGNMENT

- [Drift or gap the scorer flagged. Hook, message, proof, CTA, tone or visual.]
- [Brief pillar or section the scorer marked amber or red.]
- [Internal inconsistency between brief sections.]

Save this template for your own Loop 2 run. *Whatever both lenses flag is what your next brief must address first.*

# Reformatting



Opus Clip

**Chamomile &  
Rosemary: Real Skin  
Comfort**



**SHOULDN'T FEEL  
LIKE**

Reformatted  
using our  
engine



# How to reformat without drift

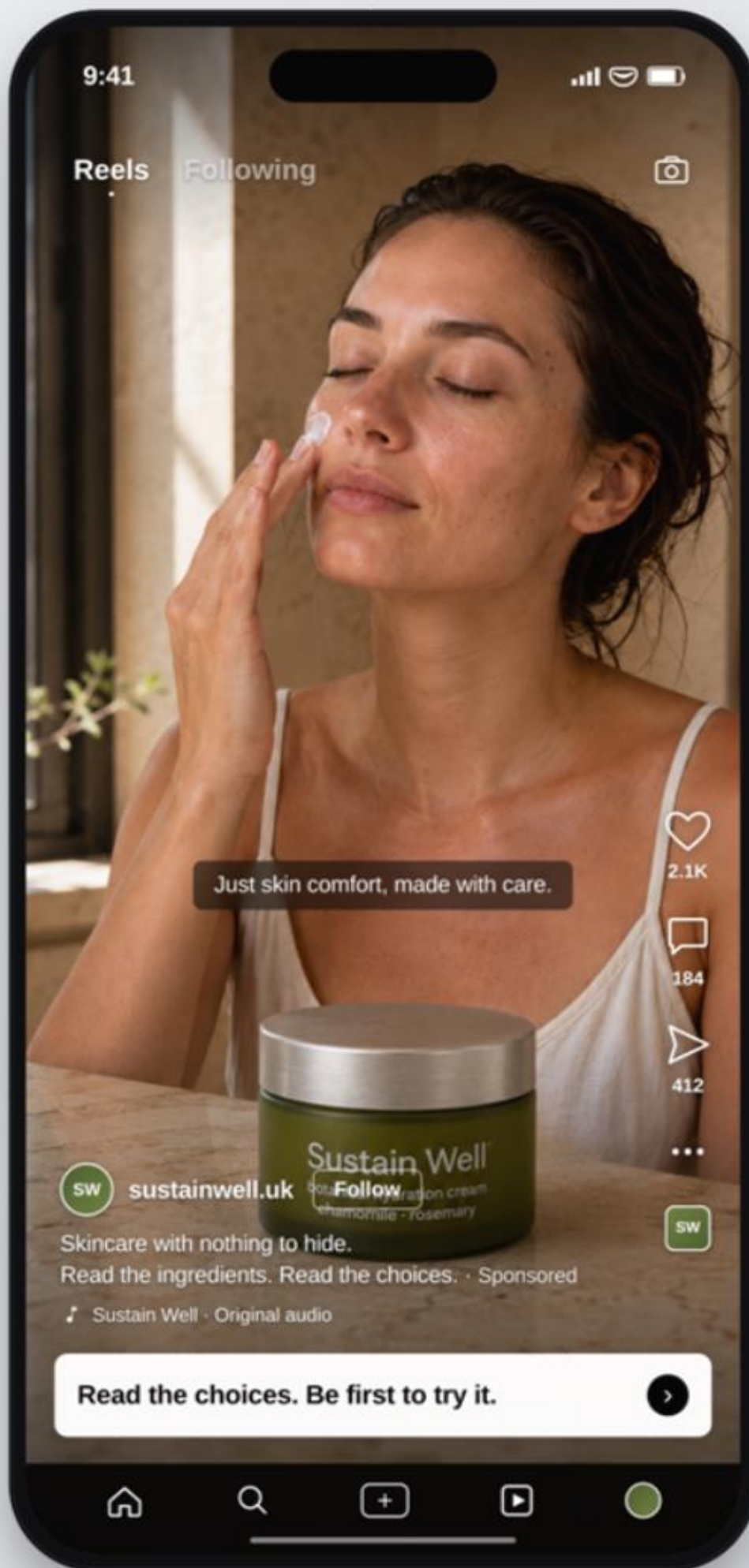
## TWO APPROCHES TO REMEMBER

- **1 Let the model write a visual identity.** A fixed description of palette, surfaces, lighting, style and exact product spec keeps every regenerated frame consistent.
- **2 Attach reference images to every prompt.** Use “match the product jar exactly to the attached reference image” so proportions and typography hold.

**The rule that prevents drift** crop in, not extend. Extending the scene adds shelves, props and background the brief never approved.

*Example failure mode: reversioning the shelf shot to vertical added a second shelf and overloaded the frame.*





# Performance Data Analysis

Test the vertical version – Loop 3

# Let's simulate the performance

Act as a senior paid social strategist. Generate synthetic but realistic Meta Ads Manager data for the attached ad, run on Instagram (Reels, Feed, Stories) over a 7-day test flight at £750.

## THE FOUR BRIEF BENCHMARKS

|                  |           |                         |           |
|------------------|-----------|-------------------------|-----------|
| Cost per sign-up | under £25 | Landing page conversion | above 4%  |
| Video hook rate  | above 25% | Time on landing page    | above 90s |

Realism rules in the prompt: a learning phase, daily variance, placement differences, day-of-week effects, and early creative fatigue. Mixed results, not a clean success story.

Meta

Ads Manager

Sustain Well · Account #842091

CAMPAIGN

Sustain Well · Waitlist pre-launch · Test flight 1

Completed

Campaigns

Ad sets

Ads

1 Jun – 7 Jun 2026

Breakdown: by day

Columns: performance

AMOUNT SPENT

£750.00

of £750 budget · 7 days

WAITLIST SIGN-UPS

22

Verified, double opt-in

COST PER SIGN-UP

£34.09

Target: under £25 · missed

VIDEO HOOK RATE

25.89%

Target: above 25% · hit

LANDING PAGE CONVERSION

4.21%

Target: above 4% · hit

FREQUENCY

1.20

Daily avg: 1.12

VERIFIED WAITLIST SIGN-UPS BY DAY

2

D1

2

D2

4

D3

5

D4

4

D5

3

D6

2

D7

Learning phase (D1–D2) · Stabilised (D3–D6) · Early fatigue (D7)

PERFORMANCE BY DAY

| DATE   | SPEND   | IMPRESSIONS | HOOK RATE | CLICKS | SIGN-UPS | CPA    | LP CR | FREQ |
|--------|---------|-------------|-----------|--------|----------|--------|-------|------|
| D1 Mon | £107.15 | 6,493       | 23.50%    | 71     | 2        | £53.58 | 3.77% | 1.05 |
| D2 Tue | £107.15 | 6,655       | 24.21%    | 79     | 2        | £53.58 | 3.28% | 1.06 |
| D3 Wed | £107.14 | 6,957       | 26.61%    | 102    | 4        | £26.79 | 4.82% | 1.10 |
| D4 Thu | £107.14 | 7,143       | 27.71%    | 114    | 5        | £21.43 | 5.32% | 1.12 |
| D5 Fri | £107.14 | 7,288       | 27.50%    | 111    | 4        | £26.79 | 4.44% | 1.14 |
| D6 Sat | £107.14 | 7,545       | 26.00%    | 94     | 3        | £35.71 | 4.05% | 1.16 |
| D7 Sun | £107.14 | 7,390       | 25.40%    | 86     | 2        | £53.57 | 2.99% | 1.18 |

Meta

Ads Manager

Sustain Well · Account #842091

CAMPAIGN

Sustain Well · Waitlist pre-launch · Test flight 1

Completed

Campaigns

Ad sets

Ads

1 Jun – 7 Jun 2026

Breakdown: by placement

Columns: performance

PLACEMENT SUMMARY

Instagram Reels

£335.50

Sign-ups13

CPA£25.81

Hook rate29.30%

Instagram Feed

£254.50

Sign-ups8

CPA£31.81

Hook rate23.60%

Instagram Stories

£160.00

Sign-ups1

CPA£160.00

Hook rate21.64%

PLACEMENT COMPARISON

Share of spend

Reels44.7%£335.50

Feed33.9%£254.50

Stories21.3%£160.00

Share of sign-ups

Reels59.1%13 sign-ups

Feed36.4%8 sign-ups

Stories4.5%1 sign-up

Cost per sign-up (lower is better, target under £25)

Reels£25.81Near target

Feed£31.81Over target

Stories£160.00Far over

PERFORMANCE BY PLACEMENT (FULL METRICS)

| PLACEMENT         | SPEND   | IMPRESSIONS | HOOK RATE | THRUPLAY RATE | CLICKS | SIGN-UPS | CPA     | LP CR | CPM    |
|-------------------|---------|-------------|-----------|---------------|--------|----------|---------|-------|--------|
| Instagram Reels   | £335.50 | 22,979      | 29.30%    | 11.60%        | 337    | 13       | £25.81  | 4.89% | £14.60 |
| Instagram Feed    | £254.50 | 17,552      | 23.60%    | 7.70%         | 221    | 8        | £31.81  | 4.57% | £14.50 |
| Instagram Stories | £160.00 | 8,940       | 21.64%    | 4.88%         | 99     | 1        | £160.00 | 1.23% | £17.90 |

# The performance analyst

## LOOP 3 · THE TOOLKIT

*ChatGPT or Claude · inputs: campaign data + campaign guidance summary + ad transcript · output: a structured performance read*

### ANALYSE THE DATA ACROSS FIVE DIMENSIONS

#### 1. Cost efficiency

Is cost per conversion within the brief's benchmark band, and where does it leak?

#### 2. Hook performance

Does the hook rate clear target, and does the daily pattern show fatigue?

#### 3. Placement mix

Which placements drove volume at what cost. Where the outliers sit.

#### 4. Day patterns

How metrics evolved across the flight. Learning phase, peak, decay.

#### 5. Creative health

Is the ad earning attention at a cost each placement charges for it?

**The output splits two ways:** short-loop fixes for this campaign now, and long-loop fixes to the brief and prompt stack for next time.

# The performance analyst

Claude, Opus 4.8 or ChatGPT 5.5



## 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.

## 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]."

# Feeding the insight back into the engine

## LOOP 3 · THE OUTPUT

*Two horizons, both fed by the performance analyst above. Most data findings close gaps in this campaign now. Most structural findings sharpen the brief for the next.*

### Improve this campaign now

#### SHORT LOOP · NEXT FLIGHT · DATA-DRIVEN

- **Cut or cap Stories.** £160 bought one sign-up..
- **Weight budget to Reels.** The only placement near target, drove 59% of sign-ups.
- **Refresh creative before fatigue.** D7 hook rate fell to 25.40% from D4 peak of 27.71%.
- **Test tighter Reels audiences.** Start refinement where CPA is already closest to target.

### Improve future campaigns

#### LONG LOOP · INTO THE PROMPT STACK

- **Numbers into rules.** Cap in the Guidance Summary what the data broke. Pause any placement above 1.5x blended CPA after 48 hours.
- **Gaps into measurements.** Every brief target needs a tracking tag live before flight. Time on landing page was written, not measured.
- **Curves into triggers.** Convert fatigue patterns into in-flight refresh rules. Refresh if hook falls two days running or frequency tops 1.15.
- **Differences into direction.** Give each placement its own creative rule in Section 10. First-frame reveal for Stories, text-priority for sound-off Feed.

*The lesson move Stories money to Reels and blended cost falls from £34 to roughly £28, just at the edge of the first-flight band. The single biggest lever does not close the gap alone, the brief has to sharpen too. That is the loop.*

# How the loops feed the engine

*Each loop surfaces a different kind of gap. Each gap has a specific home in the stack. Together, they write the next version of your stack.*

01

## CTA optimisation

### SURFACES

Gaps upstream of the ad, where the brief never declared what the brand can actually deliver.

### FEEDS

Prompts 1 and 2. Research and Brief conversion readiness.

02

## AI critique

### SURFACES

Gaps between brief and execution: missing proof, softened voiceovers, drifted anchor terms.

### FEEDS

Prompts 2 and 3. Brief anchors and Video Ad Messaging fidelity.

03

## Performance analysis

### SURFACES

Gaps between what the brief targeted and what the data actually proved.

### FEEDS

Prompts 2 and 4. Brief section 10, visual direction, and Guidance Summary rules.

*The lesson the loops critique the ad, the engine gets the upgrade. That is why every version of your stack is measurably better than the last.*

# Session 3 recap

## What you built today

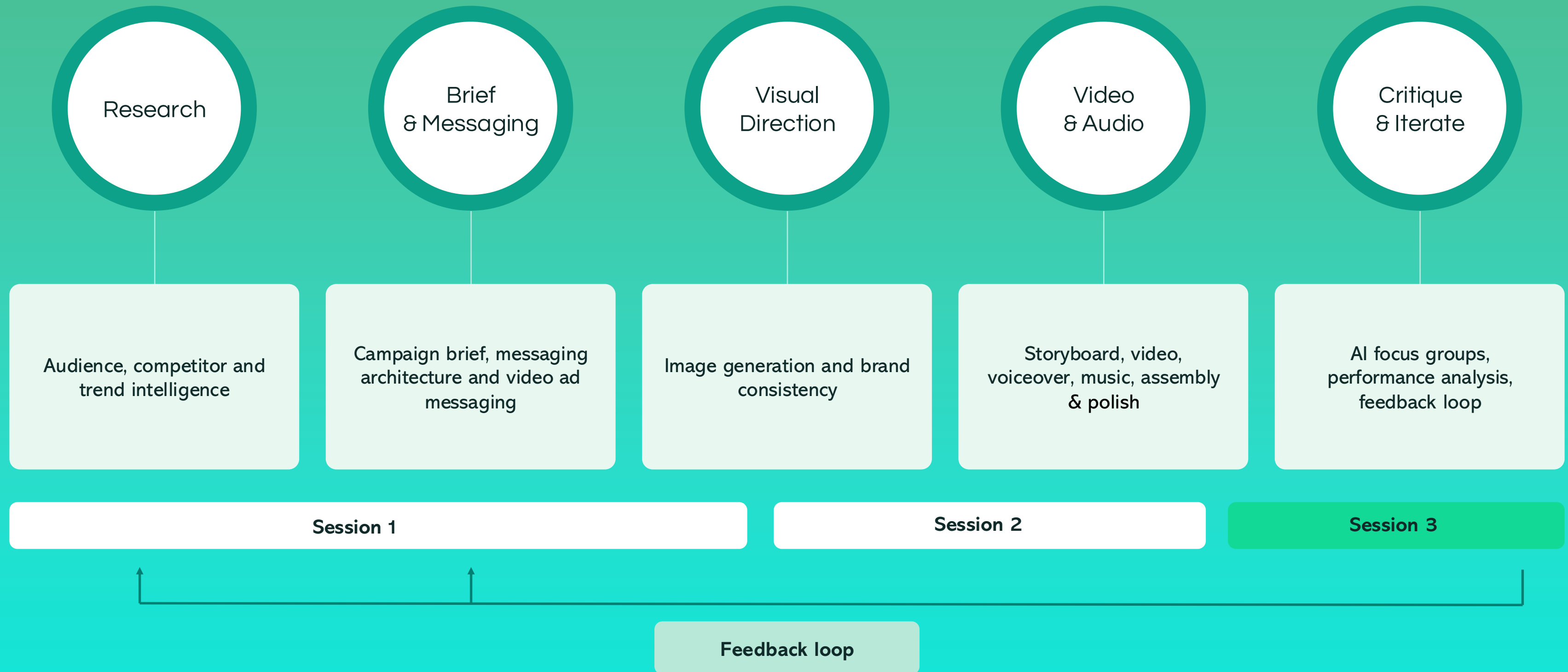
- An updated four-prompt stack, V2, with conversion readiness built in.
- A polished version of the ad with a redefined CTA and rebuilt audio.
- An AI focus group read, and a brief scorecard against your guidance summary.
- A vertical reversion of the ad, reformatted without drift.
- A performance read that turned mixed data into tangible insight that we can feed back into our engine.

## Where this sits

You now have a working campaign built across three sessions: research and strategy in Session 1, visuals, video and audio in Session 2, and in Session 3 the feedback loop that makes every campaign you run sharpen the next one.

Every campaign you run now writes an upgrade to your stack. That is the engine.

# AI workflow recap



*Every campaign you run now sharpens the next. That is the engine.*

# Copyright and AI-generated marketing content

## 1 · Output ownership

Is there enough human creativity for copyright, and in which parts? Prompts alone usually aren't enough. The UK is the outlier: it still protects computer-generated works, though a 2026 report proposes removing that.

## 2 · Input legality

Was the model trained on lawful material? Still genuinely unsettled. Don't re. 2025 US cases cut both ways on the facts. ly on "trained on the internet, therefore fine".

## 3 · Deployment risk

Does the asset create trade mark, passing off, likeness or labelling issues? Getty v Stability AI failed on copyright but partly succeeded on trade marks over synthetic watermarks.

**Practical rule for marketers** copyright is strongest where the human contribution is evident, creative and documentable. Keep your briefs, selections and edits. Avoid "in the style of [living artist]" or "make it look like Getty". For EU audiences, AI-content transparency rules apply from 2 August 2026.

# Access the feedback survey and your session resources

---

Survey



Session resources



[https://genfutures  
.co.uk/build-a-  
marketing-  
creative-studio-  
in-your-pocket-  
course-materials/](https://genfutures.co.uk/build-a-marketing-creative-studio-in-your-pocket-course-materials/)

QUESTIONS