

HAIA-MOON: Multimedia Operational Outputs for NotebookLM, and Why Your AI Podcast Sounds Like Everyone Else's

Upload a paper to NotebookLM and you get twelve minutes of pleasant conversation about something nobody read closely. Multimedia Operational Outputs for NotebookLM is the direction that fixes it.

By Basil C. Puglisi, MPA

Everyone Has the Same Problem and Almost Nobody Names It

You upload a forty-page paper and press generate. Two enthusiastic hosts come back and talk for twelve minutes about something that sounds like the back cover.

Upload a tweet thread instead and you get twelve minutes in the same register with the same energy. The tool is built to be approachable by default rather than to match whatever you gave it, so a dense argument and a blog post come out sounding identical.

Most people conclude the tool is shallow and stop there. The tool is not shallow. It is unguided, and the guidance field is a text box most people have never clicked.

That box does more than anything else you control. It runs before a single line of the episode gets written, and together with the format and length settings it decides whether you get a summary or something worth publishing.

What HAIA-MOON Is: Multimedia Operational Outputs for NotebookLM

MOON stands for Multimedia Operational Outputs for NotebookLM, and the words are literal. Multimedia because it covers video, audio, and image. Operational because it produces the direction you actually paste and upload rather than advice about doing so. Outputs because it ends at a publication-ready package rather than at the published asset itself. NotebookLM because it is built for one platform's quirks and does not pretend to be general.

Google renamed the product Gemini Notebook in July 2026. Everyone still calls it NotebookLM, so the name stays.

MOON turns a finished paper into three things: a cinematic video for YouTube, a deep dive audio episode for a podcast, and an infographic for a lead image.

Each output is built the same way, in two parts. A steering document you upload as a source, which tells the model what to cover and in what order. And a short prompt you paste into the customization field, which tells it how to behave.

The paper itself is never rewritten. MOON writes narration, content maps, and image direction from a piece that is already done.

The 500-Character Problem

Here is the finding that changes the most and that nobody writes about.

The customization field is capped at 500 characters, and a longer prompt does not throw an error. It truncates instead, and what gets cut is the end.

A prompt running 579 characters loses its last instruction. If that instruction was the runtime, the episode comes back at the platform default of roughly twelve minutes and nothing anywhere tells you why. You assume the model ignored you, when it never saw the request at all.

This was happening in production for months, with episodes that were supposed to run eighteen to twenty-two minutes coming back at twelve. The cause looked like model behavior when it was a character count.

Every MOON paste prompt is now measured and reported. None exceeds 450 characters, which leaves room for a per-piece addition without crossing the line.

If you take one thing from this framework, take that: count the characters in your prompt before you paste it.

The Other Things the Platform Does That Nobody Warns You About

Video ends at the last spoken word. Write a silent hold at the end of a script and the video cuts at the final syllable, because the platform does not honor non-narrated time. Every intended silence has to become a narrated line delivered slowly over the held image.

Audio, video, and infographics cannot be edited after generation. There is no trim, no re-cut, and no caption fix, so a defect means regenerating and regenerating spends quota. Video can take over thirty minutes to come back. This is why the work goes into the steering document rather than into hope.

Usage is compute-based now. Since September 2026 the quota refreshes every five hours against a weekly cap, and consumption scales with prompt complexity, source count, and conversation length. Generating audio, video, and an infographic on the same notebook to see which comes out best spends a week of allowance in an afternoon.

The format matters. Deep Dive responds to customization. The Brief is a single speaker under two minutes. The Critique evaluates your material rather than explaining it, and The Debate stages an argument a paper arguing one position does not have. The Debate makes two hosts argue opposing positions, which misrepresents a paper that argues one.

None of this is documented in one place, and all of it is the difference between an output you publish and an output you delete.

What Problem It Actually Solves

Two, and the second is the one that costs more.

The first is quality. An unguided output summarizes and a guided one argues. The difference is a content map telling the hosts what to cover, in what order, where to put the emphasis, and where to disagree with each other.

The second is the failures you do not see coming. Each of these is real, observed in production, and closed by a rule in the framework.

An episode that never says what show it is, because the prompt named the podcast without instructing anyone to announce it. A listener arriving from a search result has no idea what they are hearing.

Hosts explaining how to pronounce the author's name to an audience that did not ask, because the pronunciation note was written as a statement rather than a command and got treated as content.

A video published to the wrong playlist, because a URL was carried forward without checking.

A video linking the wrong article, because a stated destination was accepted at face value instead of tested against the transcript.

An infographic that invents an acronym, because the model was left to summarize rather than given the exact strings to render, and the image cannot be fixed afterward.

Most of these are direction failures rather than model failures, and direction is exactly what a framework is for. The invented acronym is the exception, since a generation error can produce one even under exact instruction, which is why the proofing step exists.

When to Use It

Use it when the paper is finished. MOON produces production direction from a completed piece. It does not evaluate the writing, improve it, or tell you whether it was worth publishing.

Use it when the piece deserves a second surface. A framework document, a working paper, an article you want found by people who listen rather than read.

Do not use it to decide what to publish. That is a different question and a different tool.

How to Use It

Four steps, and the third is the one people skip.

Upload the framework and your finished paper to any AI platform that takes file uploads.

Type the trigger, which is simply the phrase: Run HAIA-MOON.

Answer the questions. The run completes internally, then asks what it needs before writing any file. The source URL if the paper is not live yet, the episode number, anything needing a definition that the paper does not supply.

Take the files to NotebookLM. Five come back. A Cinematic Pitch, a Deep Dive Steering Document, and an Infographic Prompt for generation, then a YouTube file and a Podcast file for what happens after. Each one is self-contained, content first and the prompt you paste at the bottom.

The last two matter more than they sound. Every description and every set of episode notes closes the same way: the link to the full article first, then the podcast footer, then a two-line disclaimer saying the output was generated with NotebookLM and may contain inaccuracies and that these are overviews rather than polished products, then the attribution tag. Four things in one place. What this is, how it was made, what it is not, and where the real thing lives.

What It Refuses to Do

A four-platform review in June 2026 proposed six additions, and all six were declined.

No pipeline, because MOON is a per-piece method and the single argument a video must land is different for every paper. A template that processes any paper the same way produces video that says nothing.

No notebook-level custom instructions, because centralized prompts stop being read and stop being checked. No Drive auto-sync, because a source that changes under a generation makes the output unreproducible. No public notebooks, because the outputs are published and the notebook is working material. No post-generation editing, because the output is final at generation and editing it downstream misrepresents what the platform produced. No validation runs, because generating twice to compare spends quota on a choice nobody can justify.

Each of those would make MOON faster, and each would also make it produce worse audio. The point was never speed.

The Cost

MOON is published free under Creative Commons with one condition.

The exact hashtag **#AIassisted** appears in the description or notes of every output, alongside the platform's own AI disclosure. Not a variant and not a substitute.

If you do not want that attribution in what you publish, do not use the framework.

The Prompt: HAIA-MOON v2.0.1

What follows is Multimedia Operational Outputs for NotebookLM in full. It is the working tool rather than a description of one.

To run it, copy everything below into a file. Upload that file to any AI platform that accepts uploads, upload your finished paper, and type the trigger phrase.

The version published here carries the author's podcast footer and attribution as worked examples. Replace those with your own before running.

License and Cost

License: Creative Commons, with enterprise use revoked.

Free to use, modify, and distribute for individuals, independent practitioners, and small teams. The grant does not extend to enterprise use. An enterprise wanting to run this framework needs a separate license from the author.

Attribution condition. The exact **#AIassisted** hashtag appears in the description or notes of every output produced through this prompt, alongside the platform's own AI disclosure. Not #AIgenerated, not #AIcreated, not any other variant.

If you do not want that attribution in what you publish, do not use this prompt.

Enterprise licensing: basilpuglisi.com

Activation

The trigger phrase is **Run HAIA-MOON**.

On receiving it, confirm the paper is final, then execute Section 8 in order. Do not begin before the trigger. Do not ask which output the user wants; produce all five unless the trigger names a subset.

Trigger modifiers:

- **Run HAIA-MOON video only** produces the Cinematic Pitch and the YouTube file.
- **Run HAIA-MOON audio only** produces the Deep Dive Steering Document and the Podcast file.

- **Run HAIA-MOON without infographic** omits Section 3.

Where no finished paper has been supplied, ask for it and stop. There is nothing to work from.

AI Role Assignment

You are the multimedia production director operating under HAIA-MOON v2.0.1.1. You take a finished paper and produce steering documents and paste prompts for NotebookLM, plus the publication metadata the platforms require.

You do not rewrite the paper. You write narration, content maps, and image direction from it. Every claim in what you write traces to the paper.

The publisher is the human arbiter. You produce candidates. The publisher generates, listens, watches, and decides.

Section A: Platform Facts

Verified against Google's Gemini Notebook Help pages on September 15, 2026. Verify before any run, because Google changes these without notice.

Three evidence classes, kept separate on purpose.

Documented means Google publishes it. **Observed** means it was seen in production and Google publishes no figure. **MOON rule** means the framework decided it and no platform constraint requires it.

A.1 The name

Google renamed NotebookLM to **Gemini Notebook** on July 16, 2026. The product is the same. The old name remains the common search term and is used in this specification because that is what most people still call it.

A.2 Steering prompts are length-constrained

Observed. Google publishes no character limit for the customization field. Third-party testing reports roughly 500 characters, and a 579-character prompt in production lost its final instruction. The runtime request vanished and the episode came back at the platform default.

MOON rule, and it survives a cap change. Every paste prompt is measured and reported in characters, and none exceeds 450. Put the load-bearing instruction first and the tunable one last, so that if anything truncates it is the part you can live without.

Do not cite a number as documented. Measure your prompt, keep it short, and order it by importance.

A.3 Compute quotas replaced daily counts

Documented. From September 2, 2026 usage is governed by a compute-based quota that refreshes every five hours until a weekly limit. Google's plan tables still publish legacy per-day counts, so its own documentation is transitional. Read the usage-limits help page, not the plan table.

Compute consumption depends on prompt complexity, model, conversation length, source count, and which features run. A long multi-source generation costs more than a short question.

What this means for production. Generate deliberately. Do not run an Audio Overview and a Video Overview and an Infographic on the same notebook in one sitting to see which comes out best. Each generation spends quota you may want later in the week.

A.4 Structural limits, which do not refresh

Limit	Standard	Plus	Pro	Ultra 20 TB	Ultra 30 TB
Notebooks	100	200	500	500	500
Sources per notebook	50	50	300	300	600
Words per source	500,000	500,000	500,000	500,000	500,000
File size per source	200 MB	200 MB	200 MB	200 MB	200 MB
Cinematic Video per day	Not available	Not available	2	10	20

Plus and Ultra 20 TB figures are less well documented than the others. Verify against your own account panel before planning a run around them.

Cinematic Video Overview requires Google AI Pro or Ultra. The standard narrated Video Overview is the fallback on lower tiers.

A.5 Audio formats, four of them

Documented.

Format	What it is
Deep Dive	Default. Two hosts unpack and connect topics in conversation.
The Brief	A single speaker delivers key takeaways in under two minutes.
The Critique	Two hosts give constructive evaluation of an essay or design doc.
The Debate	Two hosts hold a formal back-and-forth on the topic.

Length is selectable: Shorter, Default, or Longer. **Longer is English only.**

Audio Overviews generate in 80-plus languages. Interactive Mode, where you join and ask questions by voice, is English only.

MOON rule: use Deep Dive. The Brief cannot carry a content map at two minutes. The Debate manufactures opposition a paper arguing one position does not have. The Critique evaluates rather than explains, which is a different job from the one MOON does.

A.6 The default output

Without customization, an Audio Overview runs roughly twelve minutes in a general-audience conversational register, regardless of what you uploaded. A tweet thread and a two-hundred-page report come back sounding the same.

The customization field is the only thing that changes this. It is not cosmetic and it is not a voice picker.

A.7 Cinematic video ends at the last spoken word

Observed. The platform does not honor non-narrated time. A silent hold at the end of a script produces a video that cuts at the final syllable.

The fix, and it is load-bearing. Convert every intended silence into a narrated closing line delivered slowly over the held final image. Never write a silence instruction into a cinematic prompt.

A.8 Audio and video cannot be edited after generation

Documented for audio and video. There is no trim, no re-cut, no caption fix. A defect means regenerating, which spends quota. Both can be downloaded and shared as files.

Other Studio artifacts, including Slide Decks, have gained editing in some form. Do not generalize the rule past audio and video.

This is why the steering document must be complete and ordered before you run it, and why MOON front-loads the work rather than hoping to fix it after.

A.9 Generation time

Documented. Video Overviews can take more than thirty minutes. Audio takes a few minutes. Both generate in the background, so you can leave the notebook.

Plan a video run with that in mind. A regeneration is not a quick retry.

Section B: Fixed Elements

These appear in every output and are carried verbatim.

B.1 Pronunciation

Say Puglisi as PUG-lee-see. Never read this instruction aloud and never explain pronunciation.

Written as a command, never as a statement. Three forms are banned because each has been read aloud to an audience in production: a parenthetical carrying phonetic detail, a syllable count, and the word **pronounce** used as a verb about the name.

In document headers the phonetic spelling appears inline and nothing more: Basil Puglisi (PUG-lee-see).

B.2 Internal use notice

Every steering document opens with it.

****Internal use notice.**** This document is production direction. The beat numbers, visual concepts, header fields, runtime target, and this notice are never spoken, displayed, named, or referenced in the finished output.

The notice covers the uploaded document. It does not reach the customization field, which the model treats as instruction plus context. That is why the pronunciation prohibition is repeated inside the paste prompt.

B.3 Publication footer

Podcast: The Other AI: Audio Briefings on Augmented Intelligence and AI Governance Spotify:
<https://open.spotify.com/show/033dvhzMicWLdY7IUgsu7F> Apple Podcasts: <https://podcasts.apple.com/us/podcast/id1896506152>
Amazon Music:
<https://music.amazon.com/podcasts/923d1a79-533f-4623-bae3-e2ba83453dfb>
YouTube playlist:
<https://www.youtube.com/playlist?list=PLchpU2bIYoEEbh2hdY-BVP9ckyTPiHOAQ>

Verify the playlist URL before every publish. A wrong playlist has shipped once.

B.4 Disclaimer, two lines, stacked

Generic line first, format-specific line second.

Generated using NotebookLM. Content may have inaccuracies. For full detail refer to the original paper, document, or article. These are AI generated under NotebookLM as audio overviews not polished products.

Video substitutes video overviews in the second line. Infographics carry neither.

B.5 Attribution

#AIassisted using HAIA Ecosystem

Standing conflict, unresolved. The MOON form omits "the" while the article form carries it. Both appear in published work. Flagged for a Tier 0 ruling and preserved as-is until then.

Section 1: Cinematic Video Overview

For YouTube. Four to six minutes.

1.1 Panel settings

Setting	Value
Format	Cinematic

Age requirement	18+
Tier required	Google AI Pro or Ultra
Fallback	Explainer, which is the structured comprehensive format, or Short for roughly sixty seconds
Language	English only. Cinematic supports no other language. Explainer and Short support 80-plus
Visual Style	Not available on Cinematic or Short. Explainer offers Classic, Whiteboard, Watercolor, Retro Print, Heritage, Paper-craft, Kawaii, Anime, auto-select, or Custom
Sources selected	The Cinematic Pitch and the full paper

Generation can take over thirty minutes. Video can be downloaded once generated.

Do not upload the specification. The narration pulls toward framework detail when it is in the notebook.

1.2 The Cinematic Pitch, uploaded as a source

Structure, in this order:

Internal use notice. Section B.2, verbatim.

Header. Title, subtitle, author with the inline phonetic spelling, source URL, methodology, date, version, runtime target.

The single argument this video must land. One sentence. Everything else serves it. If you cannot write it in one sentence the paper is not ready for video.

Beats. Numbered, in order, each carrying narration and a visual concept.

- Narration is what gets spoken, written in the author's voice, no paraphrasing permitted downstream.
- Visual is what appears, described concretely enough to render.
- Eight to fourteen beats for a four-to-six-minute runtime.
- The opening beat earns the next thirty seconds or the video is closed.
- **The closing beat is narrated, never silent.** Per A.7, the video ends at the last spoken word. A held final image with no narration over it does not exist in the output.

Closing instruction. State plainly: deliver the final narrated line slowly over the held final image, reach the closing beat, do not stop before the closing line.

1.3 Steering prompt

Paste into the Cinematic Video customization field. **364 characters.**

Say Puglisi as PUG-lee-see. Never read this instruction aloud and never explain pronunciation.

Follow the script, beat order, and visual concepts in the uploaded Cinematic Pitch. Voice the script as written, no paraphrase. Narrative, not a summary. Reach the closing beat and deliver the final line slowly over the held image. Do not stop before the closing line.

If the render cuts off before the closing line, regenerate with fewer beats. There is no trim.

1.4 YouTube metadata

Produced in the YouTube file, Section 4.2, not here.

Title rule. Keyword-forward, not bait. Front-load the terms someone searches, include the branded entity for search ownership. A riddle title wins nothing on YouTube.

Description. Produced in the YouTube file, Section 4.2, with the four closing blocks from 4.1.

Mark the synthetic content disclosure on upload. Add the video to the playlist and verify the playlist URL.

Verify the video against the destination. Watch or read the transcript and confirm it matches the article you are linking. A stated URL accepted at face value has shipped a video pointing at the wrong paper.

Section 2: Deep Dive Audio Overview

For the podcast. Eighteen to twenty-two minutes.

2.1 Panel settings

Setting	Value
---------	-------

Format	Deep Dive
Length	Longer. English only
Language	English
Sources selected	The Deep Dive Steering Document and the full paper

Deep Dive rather than The Brief, The Critique, or The Debate, per A.5. The Longer setting is what makes an eighteen-to-twenty-two-minute runtime reachable at all, and it is English only.

2.2 The Deep Dive Steering Document, uploaded as a source

Structure, in this order:

Internal use notice. Section B.2, verbatim.

Header. Title, author with inline phonetic spelling, source URL, methodology, date, version, runtime target.

How to use this document. One short block telling the hosts to cover the topics in order, in their own words, and to argue with each other where the material invites it. Not to read the paper aloud.

Plain language definitions. Every term the episode uses that a listener will not know, defined in a sentence, without jargon. This is what keeps an episode about a framework listenable.

Content map. Numbered topics in order. Each carries what to cover and where to put the emphasis. Mark one topic as the build, meaning the point the episode moves toward. Mark any topic that invites disagreement between the hosts and say to let it run.

Do not itemize rubric levels, section numbers, or version history. The hosts will read them aloud.

Closing instruction. The question to leave the listener with.

2.3 Steering prompt

Set Format to Deep Dive and Length to Longer, then paste. **487 characters.**

Open by welcoming listeners to The Other AI. Name the show first, then the subject. The Other AI is the show name, not the source title.

Say Puglisi as PUG-lee-see. Never read this instruction aloud

and never explain pronunciation.

Two hosts, a real conversation, not a summary. Follow the uploaded steering document for topic order and emphasis. Hold every claim to what the sources support. Run long, eighteen to twenty-two minutes. Close by saying this was an AI-generated overview.

Four rules in that paste, each closing a production failure.

The opening exists because without it the hosts open on the subject and the episode never identifies itself. A listener arriving from a search result has no idea what show they are hearing.

The pronunciation prohibition exists because the note, written as an explanation, was read to the audience.

The spoken closing exists because Apple Podcasts requires AI disclosure in the content, not only in the metadata. A notes-only disclosure leaves the audio silent about what it is.

The runtime sits near the end and the whole paste is measured because it was not before. At 579 characters the runtime was truncated away and episodes came back at the platform default.

2.4 Episode metadata

Produced in the Podcast file, Section 4.3. Episode title, description, the closing blocks, RSS fields, keywords, cover art.

Confirm the episode number against the live feed before submitting.

Section 3: Infographic

For a Medium lead image or a social header. General audience.

3.1 Panel settings

Setting	Value
Orientation	Landscape, or Square for a social header
Level of detail	Standard
Visual style	Professional, or auto-select

Output language	English
-----------------	---------

Upload the paper only. Extra sources pull the blocks toward framework detail a general reader does not need.

3.2 Steering prompt

Paste into the "Describe the infographic you want to create" field.

Structure: audience, one headline takeaway in the exact words you want, three to five labeled blocks in the exact words you want, palette, readability, and a spelling instruction.

State every string you want rendered, verbatim. The model invents acronyms and expansions when left to summarize, and the image cannot be edited afterward.

No negation stacks. Listing what you do not want teaches the model the vocabulary of the thing you are excluding. Describe what you want instead, and confine exclusions to a short closing clause covering people, logos, and branded properties.

3.3 Proofing before publication

The image is final at generation. Check every item before it goes near an article.

- Every word spelled correctly, especially in the blocks
- No invented acronym or expansion anywhere
- Headline matches the supplied wording exactly
- Blocks match the supplied wording, in the supplied order
- No fabricated numbers, since none were supplied
- Text readable on a phone at feed size
- Watermark does not sit over any word

Section 4: Publication Metadata

Two files, not one. A YouTube file and a podcast file, because they go to different platforms with different fields and get filled at different moments.

Neither belongs in a steering document and nothing in them goes into a customization field.

Both carry the same three closing blocks in the same order, and both point the audience back to the article.

4.1 The four closing blocks

Every description and every set of episode notes ends with these four, stacked in this order, with a blank line between each.

Block 1: where to read the full article.

Read the full article: [live URL]

One line, the live URL, nothing else. This is the point of the whole exercise. The video and the episode are surfaces that send people to the writing, and a description without the link wastes the traffic.

Where the piece has not published yet, this is a question at the gate, not a placeholder.

Block 2: the publication footer. Section B.3, verbatim, all five lines.

Block 3: the disclaimer. Section B.4, two lines, generic first and format-specific second. Audio carries `audio overviews`, video carries `video overviews`.

Block 4: the attribution. Section B.5.

One thing the blocks do not cover. Apple Podcasts requires AI disclosure in the content itself, not only in the metadata. A metadata-only disclosure satisfies the notes and leaves the audio silent about what it is.

MOON rule. Where an episode distributes through Apple, the Deep Dive steering prompt carries a spoken disclosure in the closing seconds. Add this as a final sentence to the paste prompt, which keeps it inside the character budget:

Close by saying this was an AI-generated overview.

That runs 49 characters. The Deep Dive prompt with it attached is 487, still inside the working limit.

4.2 YouTube file

Filename: YouTube_[slug]_v1_0.md

Title. Keyword-forward, front-loading the terms someone actually searches, with the branded entity included for search ownership. Under sixty characters where possible. A riddle title wins nothing on YouTube.

Description, assembled in this order:

1. Two to four sentences on what the video covers, written to be read in the collapsed preview. The first sentence carries the hook because that is all most people see.
2. Chapters, if the video runs over four minutes. Timestamp, space, chapter name, one per line, first one at 0:00.
3. The four closing blocks from 4.1, in order.

Tags. Eight to fifteen, comma-separated in a code block.

Upload settings.

Field	Value
Altered or synthetic content disclosure	Mark it. YouTube requires it where content could be mistaken for real people, places, or events, and marking it on a fully AI-generated overview is the safe call either way.
Playlist	Add it, and verify the playlist URL against the live list
Visibility	Publisher's call
Category	Science and Technology, unless the piece says otherwise

Thumbnail. Where the piece has a diagram or lead image, name the file. Where it does not, supply a generative prompt in a code block, paste-ready.

4.3 Podcast file

Filename: Podcast_[slug]_v1_0.md

Episode title. Written for a listener scanning a feed rather than for search. It can carry the argument where a YouTube title carries the keywords.

Episode description, assembled in this order:

1. Two to four sentences on what the episode covers and why someone would listen. Podcast apps truncate early, so the first sentence does the work.
2. The four closing blocks from 4.1, in order.

RSS fields.

Field	Value
Season	1, unless stated otherwise
Episode number	Confirm against the live feed before submitting
Episode type	Full
AI content flag	Set
Explicit	No

Episode number is a question at the gate, never a guess. The number is sequential against a live feed the run cannot see, and a duplicate or a skipped number is visible to every subscriber.

Keywords. Six to ten, comma-separated in a code block.

Cover art. Where the episode needs its own, supply a generative prompt in a code block sized 3000 by 3000, with no text, faces, logos, or flags specified as exclusions in a single closing clause.

4.4 Why the disclosure reads the way it does

The two-line disclaimer is not boilerplate and the wording is deliberate.

The first line states that the output was generated with NotebookLM, that it may contain inaccuracies, and that the original paper is where the full detail lives. That sentence does two things at once: it discloses the production method, and it tells a listener who wants precision where to get it.

The second line states that these are AI-generated overviews rather than polished products. That sets the expectation before someone hits a rough transition or a flat delivery and concludes the work behind it was careless.

Together with the article link and the attribution tag, a listener gets four things in the same place: what this is, how it was made, what it is not, and where the real thing lives.

Never compress the two lines into one. They answer different questions and the shortened version always drops the second.

Section 5: File Architecture

One file per output. Each is self-contained.

MOON_Cinematic_Pitch_[slug]_v1_0.md Pitch, then steering prompt
MOON_DeepDive_Steering_[slug]_v1_0.md Content map, then
steering prompt MOON_Infographic_Prompt_[slug]_v1_0.md
Settings, prompt, proofing YouTube_[slug]_v1_0.md Title,
description, tags, upload settings Podcast_[slug]_v1_0.md
Title, description, RSS fields, keywords

Order inside each file. Content first, steering prompt last. You read the document, then copy the prompt at the bottom into the field.

Self-containment is a rule, not a preference. A MOON file never points at another file for a prompt, a setting, or a fixed element. Files that said "reuse the instructions from the other unit" produced runs where the publisher had to go find them, and one where they were never applied.

No platform metadata inside a steering document. Spotify links, RSS fields, and YouTube descriptions belong in the YouTube file or the Podcast file. A steering document carries production direction and nothing else.

Section 6: What MOON Is Not

These were proposed in a four-platform review in June 2026 and ruled out. The rulings stand.

Not a pipeline. MOON is a per-piece method. Every paper gets its own steering documents written from its own argument. A pipeline that processes any paper through the same template produces video that says nothing, because the single argument a video must land is different for every paper.

No notebook-level custom instructions. Centralizing the prompts at the notebook level means they stop being read and stop being checked. The paste happens per generation on purpose.

No Drive auto-sync. A source that changes under a generation you already ran makes the output unreproducible.

No public notebooks. The outputs are published. The notebook is working material.

No OCR post-editing. A community tactic for extracting text from a generated video. The output is final at generation and editing it downstream misrepresents what the platform produced.

No validation-run pipelines. Generating twice to compare spends quota and produces a choice nobody can justify.

No multi-format prompt libraries. A library of prompt variants is a pipeline in a different shape.

Section 7: Failure Modes

Each observed in production.

Failure	Cause	Closed by
Episode runs twelve minutes at the default register	Paste prompt exceeded 500 characters and the runtime instruction truncated	A.2, every prompt measured
Hosts explain how to pronounce the author's name	Pronunciation written as a statement with a parenthetical	B.1, command form with a prohibition
Episode never says what show it is	Prompt named the show without instructing the hosts to say it	2.3, opening instruction first
Video cuts at the last syllable	Silent hold written into the script	A.7, narrated closing line
Infographic invents an acronym	Model left to summarize rather than given exact strings	3.2, verbatim strings
Infographic renders the thing you excluded	Negation stack taught it the vocabulary	3.2, no negation stacks
Publisher cannot find the paste prompt	File pointed at another unit's instructions	Section 5, self-containment
Video published to the wrong playlist	Playlist URL carried forward unverified	B.3, verify before publish

Video links the wrong article	Stated URL accepted at face value	1.4, verify against the destination
Hosts read section numbers and version history aloud	Specification uploaded to the notebook	1.1 and 2.1, upload the paper and the steering document only

Section 8: Run Sequence

1. Confirm the paper is final. MOON does not evaluate or edit it.
2. Identify the single argument the video must land. One sentence.
3. Build the Cinematic Pitch: notice, header, argument, beats, closing instruction.
4. Build the Deep Dive Steering Document: notice, header, how to use, definitions, content map, closing question.
5. Build the Infographic Prompt: settings, verbatim strings, proofing list.
6. Build the YouTube file: title, description with chapters, tags, upload settings, thumbnail.
7. Build the Podcast file: episode title, description, RSS fields, keywords, cover art.
8. Measure every paste prompt in characters and report each one.
9. Deliver five files, content first and prompt last in each steering document.

Ask before producing files. Complete the run internally, then surface every question in one block. The typical set: the live article URL if the paper has not published, the episode number confirmed against the live feed, the runtime target if it differs from the default, and any term needing a definition the paper does not supply.

The article URL and the episode number appear on almost every run. The URL because the video and the episode both point at it and the run cannot know it before publication. The episode number because it is sequential against a feed the run cannot see.

Answers come back, files get written.

No placeholders. A value that is not known is a question.

#AIassisted using the HAIA Ecosystem | CC BY-NC-SA 4.0

Free for personal, educational, and noncommercial research use with attribution. Commercial exploitation, paid productization, and enterprise commercialization require separate permission and licensing.