

How The Agora Works

The Agora is an AI-generated political roundtable podcast. All five panelists — Arthur, Eli, Nora, Grant, and Sloane — are AI characters. We don't hide that. This page explains how the show is made, why it's built the way it is, and what we do to ensure that an AI-generated political show is honest, factual, and worth your time.

What you're listening to

When you press play on an episode, you're hearing:

- **Five synthesized voices**, each designed to sound like a specific person with a specific intellectual temperament. The voices are produced through ElevenLabs.
 - **Five distinct AI characters**, each reasoning independently rather than from a shared script. The panelists don't pool their thinking; each forms its own response to what has already been said. This is a deliberate design choice — it produces real disagreement rather than coordinated consensus.
 - **Pre-verified facts**. Before each issue is debated, the facts behind it — dates, dollar amounts, vote tallies, attributions — are researched and verified through Perplexity's Agent API, and the panelists are constrained to those verified facts.
 - **An AI "director" picking each turn**. For every beat of debate, all four panelists generate candidate responses in parallel; a separate director-role model call selects the one that produces the best television.
 - **A custom orchestration layer** written in Python that runs the whole pipeline — topic planning, fact retrieval, panelist generation, turn selection, fact-checking, and voice synthesis. No third-party agent frameworks.
-

Why we built it this way

Most political podcasts are either tribal echo chambers or anodyne both-sides exercises. We wanted something else: real arguments between coherent worldviews, conducted with rigor and wit, on a predictable schedule.

The AI format makes that possible in a way humans can't easily replicate. Five thoughtful commentators with consistent personalities cannot reliably meet twice a week to argue politics in front of microphones; the booking alone is impossible. AI panelists can. And because each panelist reasons independently, none is contaminated by the others' positions before forming its own.

The format draws on the long tradition of multi-perspective political panel shows — most directly *The McLaughlin Group* — but the production technology is entirely new.

Episode structure

Each episode runs roughly 28 to 32 minutes and follows a fixed shape:

1. **Cold open** — A brief banter moment, then Arthur teases the issues to come.
2. **Issue One** — A substantive policy, institutional, legal, or geopolitical story, framed and moderated by Arthur.
3. **Issue Two** — A second substantive story.
4. **Issue Three** — A third substantive story.
5. **Lightning round** — A fourth story, deliberately lighter or more cultural, where culture intersects with politics. Compressed framing, punchier turns.
6. **Predictions** — Each panelist gives one forecast in a fixed order.
7. **Sign-off** — Arthur closes the show.

Signature musical stings mark the transitions between segments — a structural cue for the listener and a moment of acoustic relief between debates.

How an episode is made

The production pipeline runs in roughly the following sequence:

1. **Topic selection.** A typical episode is generated in autonomous mode: a single command runs the planner, which draws on the past week of U.S. political news from Perplexity and selects the episode's stories, ranked for stakes, controversy, currency, and televisability, with a deliberate culture-meets-politics tilt for the lightning round. Recent episodes' topics are excluded so that consecutive episodes never retread the same ground. When major breaking news warrants it, the same pipeline runs in single-topic mode, allowing a special to ship within twenty-four hours.
2. **Episode plan.** The selected topics are assembled into a structured plan — episode title, per-issue framing questions, and the likely angle each panelist will argue.
3. **Fact retrieval.** For each issue, Perplexity's Agent API returns verified facts with source links. Anything it can't verify is cut.

4. **The debate.** Arthur introduces each issue, names a panelist, and asks a framing question. Then, for each beat: all four panelists generate candidate responses in parallel, each working from the running transcript and the verified facts for the issue; a director-role model call selects the response that makes for the best television; and when a runner-up is strong enough, it breaks in as an interruption — the source of the show's crosstalk.
 5. **Predictions.** Each panelist generates one forecast.
 6. **Fact-check.** A model equipped with live web search reviews the complete transcript, verifies any claims that didn't come from the pre-verified set — attributions, adjacent statistics, dates — and applies surgical corrections. A separate pass verifies pop-cultural attributions.
 7. **Voice synthesis.** Each turn is rendered to audio through ElevenLabs in the panelist's designated voice.
 8. **Audio post-production.** The per-turn stems are assembled with intro and outro music beds, the opening banter, and the segment stings, then mastered to podcast loudness standards.
 9. **Quality assurance.** A human listens to the full episode and flags anything — audio glitches, a fact that slipped through, a generation artifact — for correction or a re-roll.
 10. **Show notes and transcript.** Notes and a full transcript are prepared, with the source links preserved.
-

How we built the characters

The single biggest concern people have when they hear "five AI panelists" is that the panelists will all sound alike. A great deal of work has gone into making sure they don't.

Each panelist is an extensively designed character — ideology, biography, intellectual influences, rhetorical style, interpersonal dynamics with the others, and a distinct **signature rhetorical hook**:

- **Eli** reaches for Yogi-isms — actual Yogi Berra lines and Yogi-style malapropisms — to puncture solemnity from the left.
- **Nora** reaches for specific historical-moment sports analogies (Buckner, the Music City Miracle, a Joe Montana drive) when she wants a structural point to land.
- **Sloane** reaches for hip-hop and pop music — favoring short rhyming verses — which is part of what makes her a distinctive conservative voice rather than a familiar one.

- **Grant** reaches for classic film quotes and full-scene movie analogies (*Casablanca*, *Cool Hand Luke*, *Sophie's Choice*).

These hooks aren't decoration. They are scoped, ruled, and tracked.

Lane discipline. Each signature belongs to one panelist. Eli cannot deploy movie quotes; that's Grant's lane. Sloane cannot reach for Yogi-isms; that's Eli's. This keeps the characters distinct over time rather than letting their rhetorical habits homogenize.

Earned vs. forced. Every signature deployment has to serve the rhetorical point, not merely decorate it. The standing instruction is to skip the reference rather than force it; a panelist going a whole episode without deploying their signature is preferred over a shoehorned one.

No repeats. A rolling tracker prevents the same Yogi-ism, lyric, sports moment, or film from recurring across recent episodes — so the show doesn't develop verbal tics. For Grant and Sloane this applies at the level of the whole work, not just the specific quote: if Grant leaned on *Casablanca* recently, the entire film is off the table for a stretch.

The result is panelists who feel like they have memories and habits even though each is, mechanically, a fresh and independent reasoning step. The continuity isn't in the model; it's in the structure we build around it.

How we verify facts

The single biggest risk in an AI-generated political show is fabrication — a model confidently asserting something untrue. Our defense has two layers.

Before the panel speaks: verified facts. The specific facts behind each issue are researched and verified in advance, and the panelists are constrained to them. This closes the most common failure mode of AI-generated political content: confidently inventing a vote count or a dollar figure because the model's prior is "political dialogue contains specific numbers, so insert one here."

After the panel speaks: surgical fact-checking. Panelists may still make adjacent factual claims — a statistic on an unrelated bill, a quote attributed to a Supreme Court justice, the name of a Cabinet official. After generation, a model equipped with live web search reviews the complete transcript, verifies these claims against current sources, and applies targeted corrections. It typically finds one to three per episode.

Neither layer is perfect. AI models occasionally drift from their instructions, and edge cases happen. When we catch an error after release, we correct it in the show notes and, where it matters, in re-released audio — openly.

Why an AI "director"

The most distinctive piece of the architecture is the director: a separate model call whose only job is to choose which panelist's candidate response actually airs for each beat.

We built it this way because the alternative — deciding in advance who speaks next and asking them to — produces conversations that feel scheduled rather than alive. Giving every panelist a swing on every beat, and letting a director pick the swing that lands, is what produces genuine momentum and the show's crosstalk. The director chooses for the sharpest exchange, not the "most correct" take — a distinction that matters, and one we've built into how it reasons.

What's AI and what's human

COMPONENT	MADE BY
Panelist voices (audio)	AI (ElevenLabs synthesis from human-designed voice profiles)
Panelist personalities, biographies, rhetorical styles	Human-designed
Topic selection	Human-supervised (autonomous planner uses a human-designed approach)
Fact research and verification	AI-assisted, human-reviewed
The arguments and dialogue you hear	AI-generated
Per-beat turn selection	AI (the director)
Post-generation fact verification	AI with live web search, human-reviewed
Episode structure and pacing	Human-designed
Audio mixing and mastering	Human
Show notes, transcripts, corrections	Human
Cover art and character illustrations	Human-directed, human-curated

What The Agora is not

- **Not a substitute for journalism.** We do not break news. We comment on it. Our fact pipeline relies on the work of professional reporters and primary sources.
- **Not a centrist project.** The panel spans from genuinely progressive to genuinely conservative. Centrist is one chair at the table, not the table's verdict.
- **Not impersonations.** The panelists are original characters, not modeled on specific real people; any resemblance to specific commentators is incidental.
- **Not autonomous.** No model decides what The Agora discusses or how. Topic selection, fact verification, audio production, and editorial judgment remain human-supervised. The director picks turns; it does not pick topics or set editorial direction.

Questions, corrections, criticism

We take all three seriously. Contact: hello@theagora.fm