

Demo Script & Talk Track

Museum AR Smart Glass Guide — live demonstration

HumanityAI'd — July 2026

CONFIDENTIAL — INTERNAL SALES ENABLEMENT

0. Before the meeting (setup checklist)

- [] Glasses charged, paired to the demo server; server on the demo Wi-Fi with a static IP.
- [] Demo museum loaded (the 10-exhibit office instance, or a client-relevant set if pre-built).
- [] Confirm 8-language narrations approved and serving; test one recognition live.
- [] Admin panel open in a browser tab; analytics dashboard ready.
- [] Have a printed exhibit photo or two on hand as a backup recognition target.

1. Open (2 min) — the problem

"Every flagship museum here serves visitors from a hundred-plus nationalities — but the audio guide speaks two or three languages, needs a keypad, and half the visitors never pick it up. You're spending on interpretation that most of your visitors can't use. Let me show you a different way — hands-free, in their own language, and it never sends a single visitor's data to the cloud."

2. The 'wow' — live recognition (3 min)

1. Put on the glasses (or hand them to the client). Select a language — pick **Arabic** if an Arabic speaker is present, or the client's home language.
2. Look at an exhibit. Within ~1.5 seconds the narration begins **in that language**, and the title appears on the lens.
3. Look away and at a second exhibit — narration switches automatically.

"No buttons. No app. No numbers. They looked, and it spoke. And that recognition ran on a box in this room — nothing left the building."

Switch languages live to prove multilingual depth: same exhibit, now in Chinese, now in French. This is the moment that lands.

3. The content story — curator self-sufficiency (4 min)

Open the admin panel.

1. Create a new exhibit live: title, artist, a photo upload.
2. Show the visual index updating instantly — "no training, no waiting."
3. Generate narration across all 8 languages; open the **approval queue** — "nothing a visitor hears is unreviewed; Arabic in particular requires explicit sign-off by your curator."
4. Point out: "Your team did that in under five minutes. No vendor ticket. No per-language recording contract."

4. The proof (2 min)

"This isn't a mock-up. We're running a full production instance today — 10 exhibits, 8 languages, 80 approved narrations, recognition at 99.5% similarity, on a single GPU. The same software you'd deploy."

Show the case study one-pager. Show the analytics dashboard — dwell time, popular exhibits, play-through.

5. The fit — why on-prem wins here (2 min)

- Data sovereignty: nothing leaves the museum — a procurement checkbox cloud vendors fail.
- Runs on one server; your IT keeps control.
- Scales gallery-by-gallery; starts with a 4-week pilot.

6. Objection handling

Objection	Response
"Glasses are uncomfortable / battery."	Lightweight (~76 g); sessions average 20–40 min; we tune capture for battery and heat; spare units in the fleet.
"Our content team is small."	We co-author the first 25 exhibits in the pilot; after that it's minutes per exhibit. Optional managed content service.
"Is the AI narration good enough?"	Every clip is human-reviewed before it airs; you can also record professional voice for hero exhibits.
"Security / who can access it?"	Fully on-prem, pseudonymous sessions, purge endpoint; production hardening (encryption, role-based access) is on the near-term roadmap and covered in our due-diligence pack.
"What about our existing audio guide contract?"	We run alongside; the pilot proves value before you touch the incumbent.
"Price?"	Pilot is a fixed, low-risk fee; full pricing scales with catalogue and fleet — let's scope it.

7. Close (1 min)

"Let's put this in one of your galleries for four weeks — your exhibits, your languages, your Wi-Fi — and measure it against numbers we agree today. If it doesn't move dwell time and visitor satisfaction, you've risked nothing. Can we pick the gallery?"

Next step: agree pilot gallery + date, schedule the site survey, name the content point-of-contact.