

Proof-of-Concept Playbook

A 4-week museum pilot of the AR Smart Glass Guide

HumanityAI'd — July 2026

CONFIDENTIAL

1. Purpose

This playbook lets a museum run a low-risk, time-boxed pilot of the AR Smart Glass Guide in one gallery, and lets both sides judge success against agreed numbers. It is deliberately small: one gallery, ~25 exhibits, a handful of glasses, four weeks.

2. What the pilot proves

1. Visitors can put on glasses and hear the right narration in their language, hands-free, in under ~1.5 seconds — with no app to download.
2. The museum's own content team can enrol exhibits and generate multilingual narration in-house in minutes.
3. The system runs entirely on a server inside the museum — no visitor data leaves the building.
4. Engagement can be measured (dwell time, play-through, tour completion).

3. Scope

In scope	Out of scope (pilot)
1 gallery, 20–30 exhibits	Whole-museum rollout
8 languages	Custom language beyond the 8
5–10 glasses units	Large rental fleet logistics
On-prem GPU server (loaned)	Permanent server procurement
Engagement analytics dashboard	Integration with ticketing/CRM

4. Hardware & environment checklist

- [] **AR glasses:** 5–10 × RayNeo X3 Pro (HumanityAI'd supplies for the pilot)
- [] **AI server:** 1 × GPU workstation/server, NVIDIA GPU **16 GB recommended** (8 GB works for recognition-only serving; 16 GB gives head-room for on-box narration generation). 32 GB RAM, 8 cores, 512 GB SSD.
- [] **Network:** dedicated Wi-Fi SSID in the pilot gallery, ≥ 2 access points, wired uplink to the server; server on a static LAN IP.
- [] **Charging/hygiene:** charging station + wipes at the issue desk.
- [] **Power & rack space** for the server near the gallery or in the comms room.

5. Timeline

Week	Activities	Owner
0 — Prep	Kick-off, success metrics signed, gallery + exhibit list chosen, site survey (Wi-Fi heat-map), server delivered & installed	Both
1 — Content	Photograph the 25 exhibits (1-3 angles each), enrol them, author/generate the 8-language text + narration, curator review & approval	Museum content team + HumanityAI'd
2 — Trial run	Staff induction (1 hr), internal walkthrough, tune recognition thresholds, fix content gaps	Both
3 — Live pilot	Real visitors with glasses at the issue desk; daily analytics review; collect visitor feedback	Museum ops
4 — Evaluate	Measure against KPIs, visitor survey, readout workshop, go/no-go for rollout	Both

6. Enrolment procedure (per exhibit, ~5-10 min)

1. Photograph the exhibit from 1-3 angles (well-lit, framed as a visitor would see it).
2. In the admin panel: create the exhibit (title, artist, period, category, inventory no.).
3. Upload the photos — the system builds the visual index automatically (no restart, no training).
4. Write or auto-draft the short/full description, fun fact, and narration script.
5. Generate narration for all 8 languages; **curator reviews and approves** each clip.
6. Done — the exhibit is live for the glasses.

Reference point: the internal pilot enrolls a new exhibit in under 5 minutes including 8-language narration.

7. Success metrics (agree & sign at kick-off)

KPI	Target	How measured
Recognition accuracy on enrolled exhibits	≥ 95% correct	Staff test pass + visitor reports
Time from gaze to narration	≤ 1.5 s (p95)	System metrics
Visitors completing ≥ 5 narrations	≥ 60% of equipped visitors	Session analytics
Dwell-time uplift on narrated exhibits	+30% vs. baseline	Analytics vs. pre-pilot baseline
Content enrolment by museum staff, unaided	≤ 10 min/exhibit	Timed during Week 1
Visitor experience rating	≥ 4/5	Exit survey
Data egress of visitor imagery	Zero	Network audit

8. Roles

- **Museum:** gallery & exhibit selection, content sign-off, Wi-Fi & power, front-desk staffing, visitor recruitment.
- **HumanityAI'd:** server & glasses, installation, staff training, content-tooling support, analytics, evaluation readout.

9. Risks & mitigations

Risk	Mitigation
Weak gallery Wi-Fi	Site survey in Week 0; dedicated SSID + extra AP
Content bottleneck	HumanityAI'd co-authors first 25 exhibits with the curator
Glasses comfort/fatigue	Short average sessions; comfort survey; spare units
GPU capacity for narration generation	Generate narration off-peak; 16 GB GPU recommended

10. Exit & conversion

At the Week-4 readout: KPI scorecard, visitor survey summary, curator feedback, and a rollout proposal (gallery-by-gallery plan, fleet sizing, licence & support quote). Pilot content and analytics carry forward into production — no rework.