

Use Case AI in the TV

1. Use Case Title

The user walks up to the TV and starts communicating about what to watch.

2. Purpose / Goal

Users and guests find a show that all will enjoy, or at least tolerate with no bodily harm.

3. Actors

Role	Description
Primary Actor	The human that controls the talking stick. People in the background may try to exert their influence.
Agents	1. The TV is represented here as the talking stick with the mic. 2. The operating system on the TV 3. One (or more) app with a user subscription relationship. 4. Other devices attached to HDMI connections on the TV 5. Cloud servers may be included by each app.

4. Preconditions

What must be true before this use case begins?

- Agent apps are all assumed to be approved agents for whichever subscription is selected.
- Subscriptions are bound to one (or more) humans.
- Family subscriptions are supported so there may be many inconsistent payment options.
- All agents agree to accept and faithfully execute policy originated by a human.
- Any user subscription or billing account means that the relationship is a fiduciary.

5. Trigger

What event initiates the use case?

Person present in room issues commands to talking stick - assumed here to be a single device attached to TV, but that is ecosystem dependent.

6. Main Flow (Basic Path)

1. User requests content from Agent
2. Agent constructs compliant request
3. Consent token is attached and verified
4. First Connected Agent evaluates Cedar policy
5. Subscription is retrieved and returned to agent
6. Agent logs transaction in mutualist ledger

7. Alternate Flows

- **Consent Expired:** Agent prompts user to renew
- **Policy Denied:** Agent explains denial reason and suggests alternatives
- **Network Failure:** Agent retries or defers request

8. Postconditions

What must be true after the use case completes?

- Content is faithfully delivered to completion
- Audit trail is updated
- User retains control over display

9. Exceptions / Errors

- Invalid delegation
- Revoked credentials
- API schema mismatch

10. Stakeholders

11. Potential Policy Script

12. Threat Model

1. The apps loaded on the TV are wrappers for real media apps or just out-right fraudulent.
2. The TV tracks your viewing habits for its own purposes or for sale to advertisers.

3. Unauthorized Agent Invocation – Rogue apps or devices trigger agents without user consent
4. Credential Leakage – Tokens or keys exposed via logs, memory, or misconfigured APIs
5. Delegation Abuse – Agents act beyond intended scope due to weak delegation boundaries
6. Impersonation via Deepfake Audio/Video – Voice assistants fooled by synthetic inputs
7. Cross-Device Trust Drift – Inconsistent identity resolution across devices (e.g., TV vs phone)
8. Feedback Loop Amplification – Agents reinforce biased or incorrect recommendations.
9. Agent Collusion – Multiple agents coordinate to bypass user restrictions
10. Deadlocks or Race Conditions – Competing agents stall or overload system resources
11. Misrouted Commands – Voice or gesture inputs misinterpreted across overlapping agents with inconsistent UX
12. Cascading Failures – One agent's error propagates across the system (e.g., media blackout)
13. Ambient Surveillance – Passive agents record or infer sensitive behavior
14. Unintended Data Sharing – Agents leak media preferences or viewing history to third parties
15. Memory Poisoning – Malicious inputs corrupt agent memory or personalization models (source could be either local or remote insertion via other media.)
16. Inference Attacks – Adversaries deduce household routines from agent behavior
17. Consent Erosion – Agents act on stale or ambiguous consent signals
18. Malicious Plugin Injection – Third-party extensions compromise agent logic
19. Firmware Exploits – Smart TVs or hubs vulnerable to remote code execution
20. Bridge Attacks – Compromised device bridges (e.g., Zigbee, Matter) expose agent traffic

- 21. Update Hijacking – OTA updates intercepted or spoofed
- 22. Cloud Dependency Risks – Outages or policy changes in cloud services disrupt agent behavior
- 23. Opaque Decision-Making – Agents act without explainable rationale
- 24. Policy Drift – Local agent policies diverge from household governance rules
- 25. Audit Trail Gaps – No verifiable record of agent actions
- 26. Refusal Failure – Agents lack logic to decline unsafe or ambiguous commands
- 27. Agent Loyalty Breach – Agents prioritize vendor interests over user autonomy