

Voice AI vs. Video Avatars

Why the automated medical interview should keep the applicant's camera on — and the AI's face off.

01

The Design Question

Carriers replacing in-person paramedical examinations with digital Tele-MER and Video-MER processes have solved a scheduling problem and inherited a consistency one. In a doctor-led video interview, the applicant's experience depends entirely on the individual examiner's manner on the day. An examiner who is rushed or clinically brusque during a sensitive medical interrogation produces anxiety, under-disclosure, and abandonment — and the applicant attributes none of it to the examiner and all of it to the insurance brand.

Conversational AI removes that variance. But adopting it raises a design question that must be settled before build: should the AI be represented by a lifelike video avatar with lip-sync, or should it remain a voice-only assistant while the applicant's own camera stays on?

The avatar is the more impressive demonstration and the weaker product: on cost, on reliability, and on the psychology of disclosure it loses independently on each count. The applicant's camera, by contrast, earns its place for reasons the AI's face cannot replicate. These are two separable decisions, and they should be decided separately.

02

The Cost and Performance Trap

Generating real-time, responsive video avatars requires substantial computational power and cloud-rendering infrastructure, and that cost recurs with every minute of every interview. It scales linearly with volume, working against the economics that justified automating the interview in the first place. An insurer that replaces human examiners to reduce cost-per-interview, then reintroduces per-minute rendering cost, has bought consistency at a price it did not intend to pay.

The reliability penalty is the more serious of the two. Video streaming demands bandwidth that many applicants do not reliably have; on a weak 4G or 5G connection an avatar will lag, freeze, or fall out of lip-sync at unpredictable moments during a formal medical interview. Natural conversation runs on inter-turn gaps of roughly two hundred milliseconds, and telephony standards treat one-way delay beyond about 150 ms as the point at which quality begins to degrade. An avatar spends that budget on rendering and synchronisation before a single question has been asked.

It also breaks interruption handling. When an applicant cuts in mid-question — routine in medical interviews — audio can be stopped instantly. A rendered face cannot: it freezes mid-phoneme, or continues mouthing a sentence nobody is hearing. Voice-only degrades gracefully; avatars degrade conspicuously.

03

The Psychology of Trust

A digital human face asking personal health questions tends to trigger the psychological rejection response described as the uncanny valley. Affinity does not rise steadily with realism; it collapses as a synthetic figure approaches human likeness without reaching it, and motion amplifies the effect.

The mechanism is best described as *deception friction*. When an applicant looks at a face that mimics human blinking and micro-expression, subconscious suspicion forms: they sense they are being invited to respond socially to something that is not social. The failure is not that the avatar seems insufficiently human, but that it seems almost human — a worse outcome than seeming plainly mechanical.

There is a second, more concrete cost. If mouth movement falls marginally out of step with speech, attention transfers from the question to the artefact — and in an interview whose purpose is the accurate capture of medical information, any component competing with the question works against it.

VIDEO AVATAR

- Cloud rendering cost recurs per interview minute
- Fails visibly on weak connections; breaks barge-in
- Near-human face triggers suspicion, not rapport
- Reinstates the observer that suppresses candour
- Appearance is a permanent brand commitment

VOICE-ONLY AI

- Lightweight; no rendering infrastructure to fund
- Resilient on poor networks and modest devices
- Honest by construction — plainly a tool, not a person
- No digital gaze; preserves the ACASI comfort effect
- Bandwidth stays with the verification feed

04

Psychological Safety and the Quality of Disclosure

This is the argument that matters most to underwriting, and the one most often missed. Two decades of research on computer-assisted self-administered interviewing show that people disclose sensitive information more readily to a machine than to a person, because the machine carries none of the social cues that produce desirability bias. The effect is strongest on precisely the stigmatised topics — substance use, alcohol, mental health — where applicants most often under-report on a life or health application. It is the central reason for automating the interview at all.

A humanoid avatar works against that effect. Human-looking virtual agents can trigger an interpersonal avoidance response: when an applicant perceives a digital entity looking back at them, perceived privacy falls, guardedness rises, and under-disclosure of material risk factors becomes more likely. The avatar, introduced to make the interaction feel warmer, reinstates the social observer whose absence was the point.

Stated plainly: **the avatar re-imports the judgment that the AI interview was adopted to remove**. It is not a neutral cosmetic layer on top of a working design; it degrades the mechanism the design depends on.

05

Why the Applicant's Camera Should Stay On

The case against a synthetic face does not translate into a case against video. The applicant's camera serves purposes that have no equivalent on the AI's side, and it should remain enabled.

A passive facial match at the start of the session verifies the applicant against their government-issued identity document, closing off impersonation before a single medical question is asked. The camera also changes how the session is treated: applicants on video engage with greater formality, are less likely to multitask or step away, and answer more deliberately. Awareness of a reviewable record is itself disclosure-supporting — controlled research finds people report more truthfully when they know responses are recorded and can later be reviewed. That record becomes the file the insurer stands behind at claim, and the video-backed evidence a reinsurer prices treaty capacity around.

One scope limit belongs in the design and in the customer disclosure. The feed should serve identity verification, liveness, and the session record — **not** the inference of health, demographic, or lifestyle attributes from facial structure. Facial estimation of age, sex, or BMI is offered by several vendors in this market; it is not established to underwriting accuracy, it correlates with ancestry and so creates proxy-discrimination exposure, and it engages biometric and special-category data regimes that identity verification alone does not. Declining it is a procurement advantage, and it should be stated affirmatively to the applicant.

BOTTOM LINE

The purpose of an automated VMER is not to entertain the applicant with a digital puppet; it is to provide a fast, polite, frictionless and secure route to being insured. A voice-only assistant with the applicant's camera enabled delivers cost-efficiency, because no rendering infrastructure inflates the cost per interview; technical reliability, because a lightweight audio channel survives conditions that break a video stream; and disclosure quality, because the absence of a synthetic observer is the mechanism that makes AI interviewing outperform interviewer-led screening. The avatar adds cost, adds failure modes, and works against the disclosure effect the system exists to capture. It solves no problem the voice interface does not already solve better.

REFERENCES

1. Mori, M. (1970). The Uncanny Valley. Trans. MacDorman & Kageki (2012), IEEE Robotics & Automation Magazine, 19(2), 98–100.
2. Mathur, M. B., & Reichling, D. B. (2016). Navigating a social world with robot partners: a quantitative cartography of the Uncanny Valley. Cognition, 146, 22–32.
3. Lucas, G. M., Gratch, J., King, A., & Morency, L.-P. (2014). It's only a computer: virtual humans increase willingness to disclose. Computers in Human Behavior, 37, 94–100.
4. Turner, C. F., et al. (1998). Adolescent sexual behavior, drug use, and violence: increased reporting with computer survey technology. Science, 280(5365), 867–873.
5. Dinges, L., et al. (2023). Automated Deception Detection from Videos. arXiv: 2307.06625.
6. Stivers, T., et al. (2009). Universals and cultural variation in turn-taking in conversation. PNAS, 106(26), 10587–10592.
7. ITU-T Recommendation G.114, One-way transmission time.
8. Buolamwini, J., & Gebru, T. (2018). Gender Shades: intersectional accuracy disparities in commercial gender classification. PMLR, 81, 77–91.