

Title: Survey instrument appendix for “Experts disagree on how to fight AI disinformation, but agree that health and politics need different solutions”

Authors: Alexander Loth (1), Martin Kappes (1), Marc-Oliver Pahl (2)

Date: July 29th, 2026

Note: The material contained herein is supplementary to the article named in the title and published in the Harvard Kennedy School (HKS) Misinformation Review.

Appendix A: Survey instrument

The instrument was administered in English via Google Forms between July 3, 2025, and March 15, 2026. It is organized in nine sections plus consent, future-outlook, and attribution blocks. Several items use a modality- or strategy-keyed grid (e.g., the four threat-domain matrices each elicit 16 modality × domain ratings), which is how the 34 numbered questions yield the 54 substantive data points per respondent referenced in the main text. Items marked * were required; all others were optional. Likert-type response options are stated once per scale. The full LaTeX source of the instrument is included in the Harvard Dataverse deposit (see Data Availability section); the structured summary below preserves item wording, ordering, and response options.

A.1 Consent & data protection (1 item). GDPR notice naming the data controller (Alexander Loth, Frankfurt UAS), purpose, lawful basis (explicit consent), retention, and participant rights (withdrawal, access, rectification, erasure, complaint to the Hessian Data Protection Commissioner). *Q1** — “I confirm I am 18 years or older and I have read, understood, and agree to the consent terms outlined above.” — *Yes / No* (gating).

A.2 Section 1/9 — Screening & expert group (2 items). *Q2** Primary professional role — *single choice*: AI researcher/ML engineer; disinformation/fact-checking professional; journalist covering tech or politics; policymaker/regulator; ethicist/legal scholar; other (free text). Recoded post hoc into seven categories (see Appendix B). *Q3** Link to a recent activity (recent paper, organizational role, profile link) — *free text*.

A.3 Section 2/9 — Demographics & baseline knowledge (4 items). *Q5** Years of professional experience — *numeric*. *Q6** Region of primary professional activity — *single choice*: European Union; Europe (non-EU, incl. UK); North America (USA & Canada); East & Southeast Asia; South & Central Asia; Latin America & Caribbean; MENA; Sub-Saharan Africa; Oceania; other. *Q7** Self-rated technical understanding of Large Language Models — *1–7 (Novice–Expert)*. *Q8** Self-rated familiarity with deepfake video technology — *1–7 (Novice–Expert)*.

A.4 Section 3/9 — Perceived threats, modality only (1 item). *Q9*. “How threatening is each modality for spreading disinformation in general?” — *1–7 grid (Not a threat–Severe threat) for*: AI-generated text, images, audio (voice cloning), video (deepfakes).

A.5 Section 4/9 — Perceived threats, modality × domain matrix (4 rated items + 4 optional comments). Stem (repeated per domain): “Threats in the {DOMAIN} Domain (1 = Not a threat, 7 = Severe threat).” Domains: **Political** (Q10), **Health** (Q12), **Financial** (Q14), **Social** (Q16). Each is a 4-row × 7-column grid (text/images/audio/video). Each domain is followed by an *optional* free-text comment item (Q11, Q13, Q15, Q17): “Which specific points concerning the {DOMAIN} Domain would you like to share with us?”

A.6 Section 5/9 — Key risks, top-three (1 item + 1 optional). Q18* Select up to three risks judged most urgent — *multi-select (max 3)*: election interference via deepfake video; large-scale text spambots (astroturfing); audio-clone financial fraud/scams; synthetic medical misinformation; harassment/non-consensual deepfake porn; cyberattacks (e.g., phishing); other (free text). Q19. In 1–2 sentences, why these worry you most — *optional, free text*.

A.7 Section 6/9 — Mitigation strategies, ratings (1 item + 1 optional). Q20. “Rate the effectiveness of the following mitigation strategies” — *1–7 grid (Not effective–Very effective)* for: technical detection tools; digital watermarking/provenance standards; government regulation (e.g., AI Act, DSA); media/public literacy programs; platform enforcement policies. Q21. Additional information on mitigation strategies — *optional, free text*.

A.8 Section 7/9 — Mitigation strategies, forced ranking (1 item). Q22. Rank the five strategies above from 1 (most effective) to 5 (least effective) — *full ranking, no ties*.

A.9 Section 8/9 — Best–worst scaling (3 items). Q23* Which strategy is **most** effective? — *single choice from the same five strategies*. Q24* Which strategy is **least** effective? — *single choice from the same five strategies*. Q25. Single biggest obstacle (technical, political, or economic) blocking adoption of your top-ranked strategy — *optional, free text*.

A.10 Section 9/9 — Public-literacy tools (4 items). A short methodology note describes a representative tool, JudgeGPT (<https://github.com/aloth/JudgeGPT>), in which respondents view five short news fragments and make two slider-based judgments per fragment (origin: human vs. AI; veracity: legitimate vs. fake), with accuracy feedback and badges. Q26. Effectiveness of such public-awareness tools for mitigating AI-generated fake news — *1–7 (Not Effective–Very Effective)*. Q27* Single greatest limitation — *single choice*: easily bypassed by adversarial content; reaches only tech-savvy audiences; impact fades once novelty wears off; resource-intensive to maintain; other (free text). Q28. Main strength of such tools — *optional, free text*. Q29. One concrete feature or program that would make public-resilience tools more effective — *optional, free text*.

A.11 Future outlook & final comments (2 items). Q30. One emerging or currently underestimated risk at the AI–disinformation frontier on a five-year horizon — *optional, free text*. Q31. Final thoughts or feedback on the survey — *optional, free text*.

A.12 Attribution & final submission (4 items). Q32* Handling of contribution in the final publication — *single choice*: keep responses fully anonymous; list name and affiliation in the Acknowledgments; in addition, attribute specific quotes or ideas by name. Q33* Full name — *free text*. Q34* Affiliation — *free text*. Q35. Link to preferred professional/academic profile — *optional, free text*.