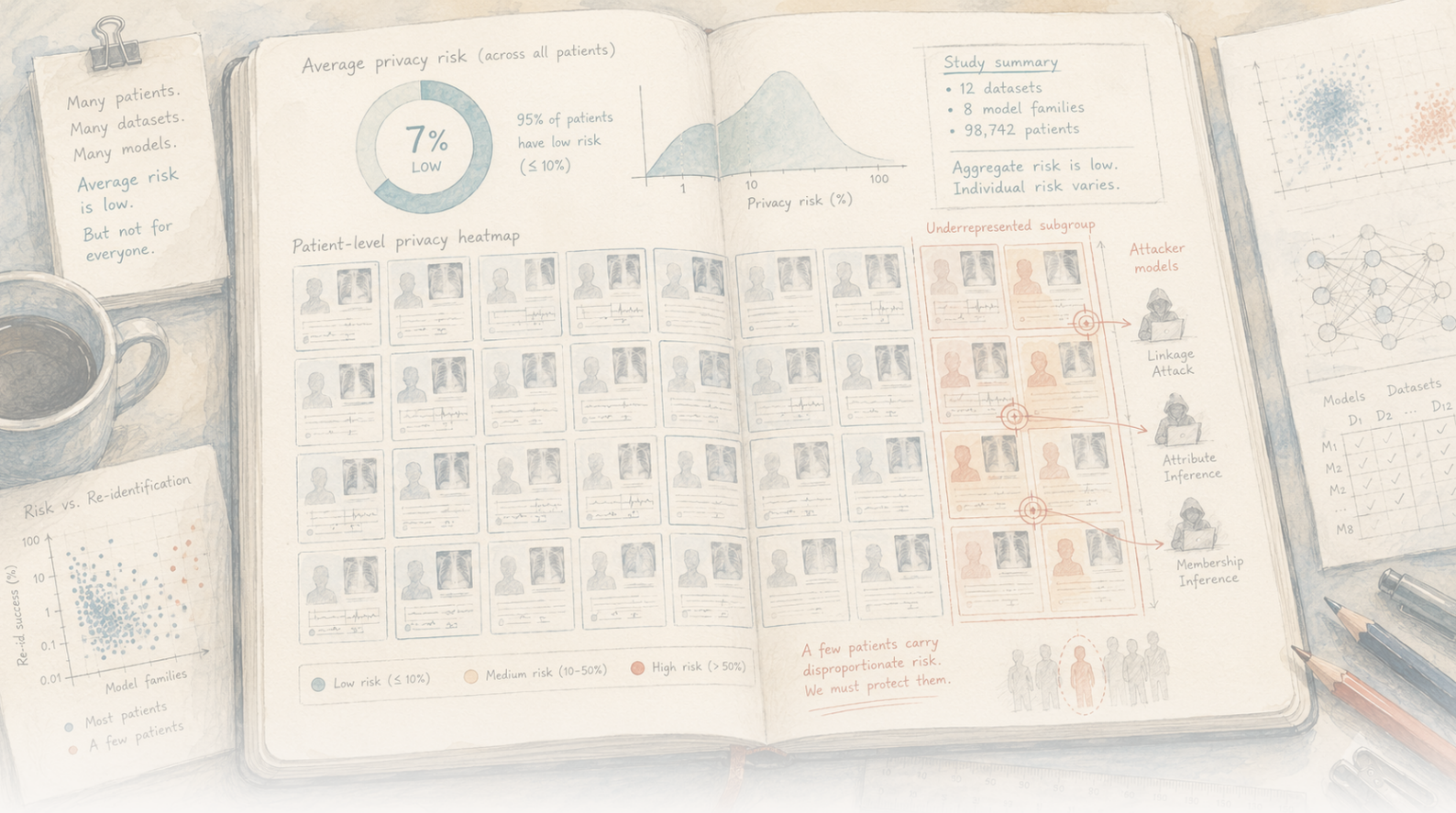




## The CLEAN PAPER

WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

# Medical AI lè jé private on average, şùgbón sì expose particular patients — underrepresented most of all

*Medical-AI model tí a pe 'privacy-preserving' sáà rest lóri average number kan. Study yíi argue number yèn wrong test. Studying membership inference attacks — which reveal whether a specific person's record was in training data and can betray sensitive disease membership — authors measured risk per patient rather than aggregate across seven medical datasets and many models. Pattern: models safe on average can let attacker identify specific individuals almost perfectly ( $AUC \geq 0.95$ ); exposed are systematically underrepresented groups; risk worsens as models grow. Authors don't say abandon medical AI — measure privacy per patient, control access, use differential privacy. Core: 'private on average' is not privacy guarantee.*

Written by Lucio Vaglio · Cross-checked by Michele Renda · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · AI-assisted, human-reviewed

Paper: *Disparate privacy risks from medical AI*, Moritz Knolle, Georg Kaissis, Daniel Rückert and colleagues (Technical University of Munich; Imperial College London; Hasso Plattner Institute), Nature (2026) · DOI: 10.1038/s41586-026-10688-0

# Privacy guarantee jẹ ilérí sí ènìyàn kan, kì í ẹ sí average

Nígbà tí a bá pe medical-AI àwòṣe kan ní “privacy-preserving,” claim náà sáà rest lórí number kan: across gbogbo aláìsàn tí dátà wọn train àwòṣe, average chance pé membership individual kan lè be inferred low. Èyí sounds reassuring. Quiet, uncomfortable kókó iwé iwádíí yí ni pé wrong number ni.

Privacy kì í ẹ average. Ó jẹ promise sí individual kọọkan — pé being in training set kò ní come back to expose wọn. Average lè keep promise fún almost everyone, sùgbón break completely fún few. olùwádíí measure àṣírí ewu ọkan aláìsàn at a àkókò, at resolution no ọkan lò sáájú, wọn sì rí exactly that: àwòṣe tí look láiléwu in aggregate lè leak membership of specific individuals almost perfectly — individuals tí wọn leak sì disproportionately jẹ those already least protected.

## Ohun tí àwọn olùkòwé ẹ

egbé — led by Moritz Knolle, pèlú Daniel Rückert (Technical University of Munich) àti Georg Kaissis (Hasso Plattner Institute), plus colleagues at Imperial College London — gba well-known àṣírí threat kan, **membership inference ìkọlù**, wọn sì yí ibèèrè padà. Dípò “on average, how often does ìkọlù succeed across dataset?”, wọn bèèrè “fún aláìsàn yí pátó, how exposed are they?”

Wọn run per-patient itupalẹ kojá **seven established ìṣẹ̀gùn datasets**, dátà types very yàtò — ìṣẹ̀gùn imaging, electrocardiograms, electronic health àkọ̀sílẹ̀ — àti 200 àwòṣe each. Fún aláìsàn kọọkan nínú àwòṣe kọọkan, wọn ìṣíró àfọjúsùn how confidently attacker lè tell whether àkọ̀sílẹ̀ person yẹn part of training dátà, then break èsì down by egbé: disease status, self-reported race, insurance, sex, imaging protocol.

### Kí ni membership inference attack gidí?

Leak níbí subtle ju “àwòṣe spits out your àkọ̀sílẹ̀.” Ó jẹ nípa *membership* — fact pé dátà rẹ wà nínú training set.

Bèrẹ pèlú ohun tí àwòṣe bá yí normally ẹ. O fún un ní scan — chest X-ray fún example — ó sì give probabilities: 78% chance pneumonia, plus readings fún cardiomegaly, oedema, consolidation. Everyday diagnostic answer yẹn nìkan ni ìkọlù lo. Nothing exotic.

Weakness: àwòṣe sáà **sí i confident** díẹ lórí exact examples tí a train rẹ on ju ones never seen. Think of student tí secretly saw exam sáájú — on practised ibèèrè answers come too quickly, too confidently. Figuring out láti that giveaway whether particular àkọ̀sílẹ̀ was ọkan of examples àwòṣe studied ni “membership inference.”

Attack step by step: attacker wants to know whether particular person’s scan lò to train àwòṣe. Wọn **kò** ask àwòṣe fún àkọ̀sílẹ̀ — wọn already hold candidate scan (person own tàbí close copy). Wọn send it once as ordinary diagnostic request, note ìgbékẹ̀lẹ̀. Then fi wé whether ìgbékẹ̀lẹ̀ looks sí i like àwòṣe that *had* trained on scan tàbí ọkan that *hadn’t* — ifiwéra they lè make cheaply by training own stand-in (“reference àwòṣe”) to learn tell-tale over-confidence on ordinary computer. If gidí àwòṣe unusually sure nípa scan — as if recognizes it — lágbára ẹrí scan was in training dátà.

Unsettling part: request looks exactly like normal clinician query; àṣírí àmì hidden inside ordinary

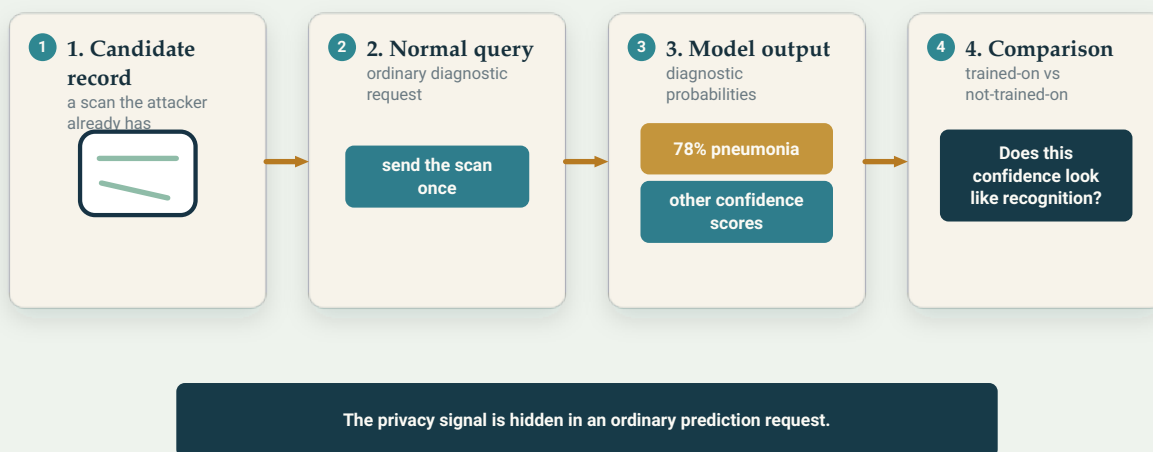
prediction. That makes ewu concrete — though demonstrated lábẹ́ stated lab assumptions, not caught in wild.

Why membership matters if àkọsílẹ̀ itself never leaks? Because *membership is a fact*. If àwòṣe trained on aláìsàn receiving particular cancer immunotherapy, confirming your àkọsílẹ̀ is in it reveals you probably had that cancer — information insurer/employer yẹ kí never infer. Content stays sealed; fact of belonging escapes.

ònkòwé ipò assumptions plainly — access to ordinary àwòṣe predictions, candidate àkọsílẹ̀, attacker's own reference àwòṣe — not as how-to, ṣùgbọ́n so threat lè be reasoned nípa rather than hand-waved.

## A membership attack can hide inside a normal prediction

*The attacker does not ask for the record. They already hold a candidate and query the model once.*



*Mechanism only: no pseudocode, no attack threshold, no recipe.*

Attack kò need special àṣírí API. Normal diagnostic query lè leak membership àmì if àwòṣe unusually confident on àkọsílẹ̀ it has seen ṣáájú.

Original Aurora diagram — The Clean Paper CC BY 4.0

## Ohun tí wọn rí

Findings mètá, each sharper than last.

**Averages hide exposed aláìsàn.** Measured in aggregate — usual way — many àwòṣe look reassuringly private: ìkọlù no better than chance fún most aláìsàn. Per-patient view yàtò. As Knolle put it in ìwádìí announcement, previous assessments “have nìkan ever measured the average ewu across gbogbo aláìsàn. We examined the ewu at the ìpele of individual aláìsàn fún the first àkókò — àti it paints a very yàtò picture.” Fún some individuals, ìkọlù

succeeds **almost perfectly** — ìkọlù AUC **0.95 tàbí higher** (0.5 coin toss, 1.0 flawless) — even while dataset-wide average looked no better than chance.



Average lè sit near chance while small tail of àkọsílẹ̀ remains much sí i exposed. ìwé ìwádíí kókó: àṣírí has to be checked at aláìsàn ìpele, not nìkan aggregate.

Original Aurora diagram — The Clean Paper CC BY 4.0

**Exposure unequal — lands on underrepresented.** aláìsàn most vulnerable to near-perfect ìkọlù were systematically láti egbé **underrepresented in dátà**: minority ethnicities, rare-disease phenotypes, unusual imaging characteristics. Nínú electronic-record dataset, Black aláìsàn appeared **31% sí i often** than expected among most-vulnerable àkọsílẹ̀; nínú mammography set, scans suspicious fún malignancy overrepresented in danger zone by **+1,179%**. (Figures are examples, not whole picture; kan náà skew across several egbé; they are *relative* over-representations nínú small extreme-risk tail — how much sí i often these aláìsàn appear among most-exposed àkọsílẹ̀, not fraction of Black tàbí suspicious-mammography aláìsàn exposed.) àwòṣe has fewer similar examples to blur these aláìsàn sínú, so àkọsílẹ̀ stand out — exactly what membership ìkọlù detects. Privacy ìkùnà not random; it concentrates on people already at margins.

**Bigger àwòṣe make it worse.** Number aláìsàn exposed to near-perfect ìkọlù rose sharply pẹ̀lú àwòṣe capacity — nínú dermatology dataset, share climbed láti essentially zero in smallest àwòṣe to nípa **ọ̀kan in ten** in largest. As ìṣẹ̀gùn AI àwòṣe get larger/sí i capable, òhàkòwé warn this ewu gets sí i severe, not less.

Rückert summary blunt: “èyí kì í ṣe a tolerable ewu. Health dátà is highly sensitive.”

## Kí ló dé tí “láléwu on average” fi wrong ìdánwò?

Temptation ni lati read low average ewu as clean bill of health. Core lesson ìwé ìwádíí: averaging here not merely imprecise — it measures wrong thing.

Privacy guarantee meaningful nikan if holds fún person most at ewu, not person in middle. àwòṣe where 999 in 1,000 aláìsàn unrecoverable sùgbón ọ̀kan identified almost perfectly kì í “99.9% private” in sense that matters to that ọ̀kan person — àti if that ọ̀kan person predictably rare-disease aláìsàn tàbí ethnic-minority aláìsàn, metric not just incomplete, quietly discriminatory. It reports ààbò fún majority àti calls it ààbò fún gbogbo.

Shift ìwé ìwádíí forces: láti *how private is àwòṣe on average?* to *who is most exposed aláìsàn, àti who are they?* Different ìbéèrè, nikan second is àṣírí ìbéèrè.

## Ohun tí èyí kò fi ẹ̀rí múlẹ̀ — tàbí claim

- Kò say ìṣẹ̀gùn AI yẹ kí be abandoned. ò̀nkòwé framing is mitigation, not retreat: measure àti fix ewu, don’t stop building.
- Kò mean every ìṣẹ̀gùn àwòṣe leaking, tàbí any deployed àwòṣe attacked. It fi hàn *ewu exists àti unequally distributed*, standard aggregate metrics miss it.
- Kò mean your àkọ̀sílẹ̀ already exposed. Attack needs specific conditions — àwòṣe access, candidate àkọ̀sílẹ̀, attacker infrastructure — not casual capability.
- Kò fi hàn aláìsàn *content* leaks. Membership — fact of inclusion — leaks, dangerous fún yàtò reason, not àkọ̀sílẹ̀ dump.
- Kò dín kù disparities to ọ̀kan àkọ̀lé number. Strong reproducible èsì is *àpẹ̀ẹ́* — aggregate metrics understate individual ewu, underrepresented aláìsàn bear most — across datasets àti hundreds àwòṣe.

## Báwo ni ẹ̀rí ẹ̀ lágbara tó?

Empirical methodological èsì, robust. Pattern — aggregate àṣírí metrics understate per-patient ewu, residual ewu concentrates on underrepresented ẹ̀gbé, worsens pẹ̀lú àwòṣe capacity — held across **seven datasets of yàtò dátà types** àti many àwòṣe. Breadth is what makes measurement claim credible rather than one-dataset artefact.

Two honest caveats. First, demonstration of *ewu*, measured by ìkọ̀lù ò̀nkòwé built; it characterises how well capable attacker *lè* do lábẹ̀ assumptions, not ayé gidi ìkọ̀lù frequency. Second, vivid numbers — near-perfect success fún some aláìsàn — describe worst-off individuals *by àpẹ̀ẹ́*; whole kókó, yẹ kí be read as “tail much heavier than average implies,” not “most aláìsàn exposed.”

Appropriate stance neither alarm nor dismissal: careful multi-dataset ìwádíí showing standard way we certify medical-AI àṣírí blind to worst ọ̀ràn — àti worst ọ̀ràn fall on aláìsàn pẹ̀lú least margin to spare.

## Kí nìdí tí ó fi ẹ̀şẹ̀ pàtàkì?

Two things make this sí i than technical footnote.

First, it changes what “privacy-preserving” yẹ́ kí be allowed to mean. àwòşẹ̀ lẹ̀ have genuinely low average àşírí score àti still hide subset of aláìsàn near-perfectly identifiable by membership ìkọ̀lù. Practical demand concrete: assess àşírí ewu **per aláìsàn** şáájú release, control access to deployed àwòşẹ̀, lò techniques like **differential àşírí** — small carefully calibrated ariwo added during training to blunt membership ìkọ̀lù, at gidi managed cost to àwòşẹ̀ usefulness (àşírí–utility trade-off explicit, not free). “We checked average” yẹ́ kí stop counting as checked.

Second, it braids àşírí pẹ̀lú fairness. Same ẹ̀gbẹ̀ underrepresented in ìşẹ̀gùn dátà — already served worse by ìşẹ̀gùn AI — turn out to have least àşírí protection. Field working hard on first inequity kò lẹ̀ treat second as someone else’s department. Same people.

Reassuring itàn — *we measured it, average low, we’re fine* — is itàn ìwé ìwádíí dismantles. Not to scare people láti technology, şùgbọ̀n to move standard where it belongs: a promise you lẹ̀ claim nìkan if checked fún person most likely harmed.

## Àkótán kedere

Membership inference ìkọ̀lù try determine whether specific person’s àkọ̀şílẹ̀ was in AI training dátà — membership itself lẹ̀ reveal sensitive facts even when àkọ̀şílẹ̀ never ojú. Prior ìşẹ̀ measured ìkọ̀lù success *on average*. ìwádíí yíí measured *per aláìsàn* across seven ìşẹ̀gùn datasets (imaging, ECG, electronic àkọ̀şílẹ̀) àti many àwòşẹ̀, rí aggregate metrics badly understate ewu: some individuals identified almost perfectly (ìkọ̀lù  $AUC \geq 0.95$ ) even when dataset average láílẹ̀wu; most vulnerable systematically láti underrepresented ẹ̀gbẹ̀ (minority ethnicity, rare disease, unusual imaging); ìşòro grows pẹ̀lú àwòşẹ̀ ìwọ̀n. ònńkọ̀wé do not argue abandon ìşẹ̀gùn AI — they say measure àşírí per aláìsàn, control àwòşẹ̀ access, lò differential àşírí. Takeaway: “private on average” kì í ẹ̀şẹ̀ àşírí guarantee, àti ìkùnà hit people already least protected.

## Àyèwò láìsí àşẹ̀jù

**Ohun tí ìwé ìwádíí fi hàn:** Across seven ìşẹ̀gùn datasets àti many àwòşẹ̀, per-patient membership-inference ewu much higher fún some individuals than aggregate metrics tọ́ka sí — up to near-perfect ìkọ̀lù success ( $AUC \geq 0.95$ ) — residual ewu disproportionately falls on underrepresented ẹ̀gbẹ̀ àti grows pẹ̀lú àwòşẹ̀ capacity.

**Ohun tó şẹ́ ẹ̀gbà şùgbọ̀n tí a kò fi ẹ̀rí múlẹ̀:** That ayé gidi attackers already exploit this against deployed ìşẹ̀gùn àwòşẹ̀; ìwádíí demonstrates *capability àti distribution* lábẹ̀ assumptions, not frequency in wild.

**Ohun tí kò fi hàn:** Medical AI yẹ́ kí be abandoned; every àwòşẹ̀ leaks tàbí specific deployed àwòşẹ̀ breached; aláìsàn àkọ̀şílẹ̀ *contents* rather than membership exposed; single quantified disparity figure — robust èsì is àpẹ̀rẹ̀, not òkan number.

**Main limitations:** Measures worst-case ewu via ònńkọ̀wé’ ìkọ̀lù, not observed incidents; dramatic figures describe

most-exposed individuals by àpẹrẹ; exact per-group disparities shown as consistent àpẹrẹ across datasets, not òkan number.

**Confidence wo ni gbogbogbò reader yẹ kí ó ní?** High that aggregate àşırı metrics understate individual ewu, residual ewu concentrates on underrepresented aláìsàn, worsens pẹlú àwòşe ìwòn — demonstrated across many datasets/àwòşe. Moderate on ayé gidi frequency of ìkọlù, not measured. Safe reading: not “ìşègùn AI leaks your dátà,” not “àşırı solved,” şùgbọn “standard àşırı check is blind to own worst ọràn, àti those ọràn fall on most vulnerable aláìsàn — so check gbọdò change.”

---

## Àwọn orísun

Knolle et al., Disparate privacy risks from medical AI, Nature (2026)

<https://doi.org/10.1038/s41586-026-10688-0>

Technical University of Munich — press release, 'Study reveals privacy risks in medical AI' (2026)

<https://www.tum.de/en/news-and-events/all-news/press-releases/details/study-reveals-privacy-risks-in>

Medical Xpress (2026) — coverage of the study

<https://medicalxpress.com/news/2026-06-patient-groups-vulnerable-privacy-medical.html>