



The CLEAN PAPER

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

चकित्सकीय AI औसतन नजिदख सकतीहै, फरि भी कुछ मरीजोकोउजगर कर सकतीहै — खसकर कम प्रतनिधित्व वाले लोगोको

“गोपनीयतासुरक्षति” कहे जमे वाले चकित्सकीय AI मॉडल अक्सर एक औसत संख्यापर टकिे हते हैं। यह अध्ययन कहताहै कि वह कसैटीपर्याप्त नहींहै। लेखकोने membership inference हमलोकाअध्ययन किया— ऐसे हमले जोयह अनुमम लगते हैं कि कसिव्यक्ति कारकिँड मॉडल के प्रशक्षिण डेटामें थायानही और इस तरह कसिबीमरीसे जुड़ीसंवेदनशील जमकरीकासंकेत दे सकते हैं। सप्त चकित्सकीय डेटासेट (इमेजगि, ECG और इलेक्ट्रॉनिक रकिँड) तथाअनेक मॉडलोमें उन्हेंमे जोखमि कोऔसत के बजय मरीज-दर-मरीज ममा नतीजा औसतन सुरक्षति दखिने वालामॉडल भीकुछ व्यक्तियोंकोलगभग पूरीतरह पहचमने दे सकताहै (attack AUC ≥ 0.95); सबसे अधिक उजगर लोग व्यवस्थति रूप से कम प्रतनिधित्व वाले समूहोसे आते हैं — अल्पसंख्यक जतीयता, दुर्लभ रोग याअसमम्य इमेजगि — और बडे मॉडल में समस्याबदतीहै। लेखक चकित्सकीय AI छेड़ेने कोनहीकहते; वे प्रत-मरीज गोपनीयताममने, मॉडल तक पहुँच नयित्त्रति करने और differential privacy इस्तेमल करने कीसलह देते हैं। असुवधिजनक नष्कर्ष: “औसतन नजि” हमेाकसिव्यक्ति कीगोपनीयताकी गारंटीनहीहै।

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

गोपनीयतागारंटीकसीऔसत से नहीं हर व्यक्ति से किया गया वदा है

जब किसी medical-AI मॉडल को “privacy-preserving” कहा जाता है, तो यह दवा अक्सर एक ही संख्या पर टिका होता है: जिन सभी मरीजों के डेटा से मॉडल train हुआ, उनमें औसतन किसी व्यक्ति की सदस्यता का अनुमान लगाने की संभावना कम है। यह आश्वस्त करने वाला लगता है। इस शोधपत्र का शांत लेकिन असुविधाजनक बटु है कि यह गलत संख्या है।

गोपनीयता कोई औसत नहीं है। यह हर व्यक्ति से किया गया वदा है—कि प्रशिक्षण set में हमें भविष्य में उन्हें उधार नहीं करेगा और कोई औसत लगभग सभी के लिए यह वदानभित्ति हुए कुछ लोगों के लिए इसे पूरी तरह तोड़ सकता है। शोधकर्तृओं ने गोपनीयता जोखिम को patient-by-patient, पहले कभीन इस्तेमाल की गई resolution पर मिला और ठीक यही पता समग्र में सुरक्षा दिखाने वाले मॉडल कुछ खास व्यक्तियों की सदस्यता लगभग पूरी तरह रसिद्ध कर सकते हैं—और जिनकी सदस्यता रसिद्ध होती है, वे अनुपत्ति से अधिक वही लोग हैं जिन्हें पहले से सबसे कम सुरक्षामाली हुई है।

लेखकोने क्या किया

Moritz Knolle के नेतृत्व वाली टीम—Technical विश्वविद्यालय of Munich के Daniel Rückert और Hasso Plattner Institute के Georg Kaissis, साथ में Imperial College London के सहयोगी—ने सदस्यता अनुमान हमलानस के प्रसिद्ध गोपनीयता threat को लिया और सवाल बदल दिया: “डेटासेट में औसतन यह हमला कतिनीबार सफल होता है?” पूछने की बजाय उन्होंने पूछा: “इस खास मरीज के लिए एक्सपोजर कतिना है?”

उन्होंने यह per-patient विश्लेषण सत् स्थिति चकित्सकीय डेटासेट पर चलाया जिनमें बहुत अलग डेटा types थे—चकित्सकीय इमेजिंग, electrocardiograms और इलेक्ट्रॉनिक स्वस्थ रिकॉर्ड—और हर डेटासेट पर 200 मॉडल तक। हर मॉडल में हर मरीज के लिए उन्होंने अनुमान लगाया कि attacker कतिने विश्वास से बता सकता है कि उस व्यक्ति का रिकॉर्ड प्रशिक्षण डेटा में था या नहीं फरि परिणाम को बीमारी स्थिति, self-reported race, insurance, sex और इमेजिंग प्रोटोकॉल जैसे समूह में बाँटा।

Membership inference attack वास्तव में क्या है

यह रसिद्ध “मॉडल आपका पूरा रिकॉर्ड बहर उगल देता है” जैसी चीज नहीं है। यह सदस्यता के बारे में है—सिर्फ यह तथ्य कि आपका डेटा प्रशिक्षण set में था।

पहले समझें कि ऐसा मॉडल समझ्यतः क्या करता है। आप उसे कोई स्कैन—मम लें chest X-ray—देते हैं और वह probabilities लैटैस है: pneumonia की 78% संभावना, साथ में cardiomegaly, oedema, consolidation जैसी readings। हमला इसी रोज़मर्रा के नैदानिक उत्तर का उपयोग करता है। कुछ असंग्रहण नहीं।

कमजोरी यह है कि मॉडल अक्सर उन सटीक उदहरण पर थोड़ा अधिक confident होता है जिन पर वह train हुआ था। बनसिबत उन उदहरण के जिन्हें उसने कभी नहीं देखा। इसे ऐसे student जैसा समझें जसिने चुपके से exam पहले देख लिया हो—practice किए प्रश्न पर उत्तर थोड़ा ज़्यादा तेजी और विश्वास से आते हैं। इसी संकेत से यह पता लगाना कि कोई खास रिकॉर्ड प्रशिक्षण उदहरण में था या नहीं “सदस्यता अनुमान” है।

हमलाचरण by चरण ऐसा है। कोई व्यक्ति जन्मा चाहता है कि किसी खास मरीज का स्कैन मॉडल प्रशिक्षण में इस्तेमाल हुआ था या नहीं वह मॉडल से रिकॉर्ड माँगता नहीं—उसके पास उम्मीदवार स्कैन पहले से है, मरीज का अपना या बहुत करीबी प्रतिलिपि। वह स्कैन को समझ्य नैदानिक request की तरह एक बार भेजता है और उत्तर का विश्वास टपिपणी करता है। फरि वह देखता है कि यह विश्वास उस मॉडल जैसा लगता है जसिने स्कैन पर train किया था या उस

जैसा जसिने नहीं किया था—यह तुलनावह अपने computer पर एक stand-in “संदर्भ मॉडल” train करके सस्ते में सीख सकता है, बनाविशेष hardware के। यदि वस्तुवकि मॉडल इस स्कैन पर असमम्य रूप से sure है—ममोउसे पहचमताहो—तोयह मजबूत सक्ष्य है कि स्कैन प्रशक्षिण डेटामें था।

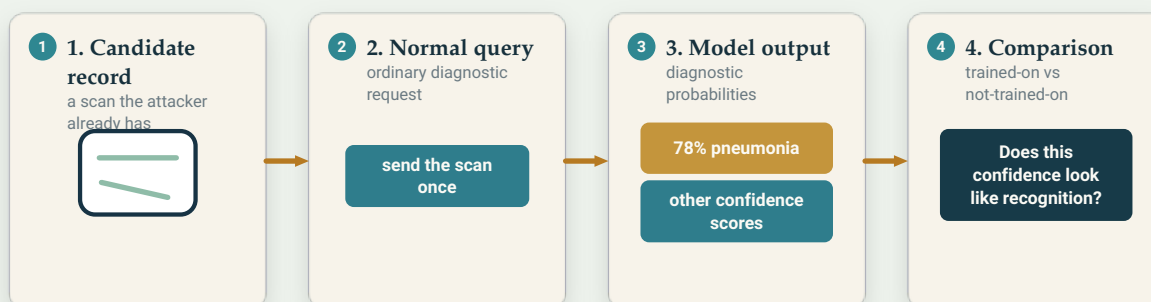
असुवधिजनक बत यह है कि request हमलाजैसीदखितीहीनही यह वहीनैदमकि प्रश्न है जोचकित्सक कर सकता है, और गोपनीयतासंकेत समम्य पूर्वमुमम के भीतर छुपिहै। इसीकारण जोखमि ठोस है—हलँकअभीतक इसे बतई गई प्रयोगेशलामम्यताएँ के तहत दखियागयाहै, वस्तुवकि दुनियामें पकड़ानहीगया।

यदरिक्छिड खुद रसिम नहीहेतातोसदस्यताक्योमयने रखतीहै? क्योकि सदस्यतास्वयं एक तथ्य है। यदिमॉडल खस cancer immunotherapy पाए मरीज पर train हुआ, तोयह confirm करनाकि आपकारक्छिड उसमें है, यह बतासकताहै कि आपकोशयद वहीcancer था—ऐसीबत जिसकाअनुमम insurer याemployer कोकभीनहीलगापता चहिए। Content बंद रहताहै; बहर नकिलताहै belonging कातथ्य।

लेखक मम्यताएँ सफ लखिते हैं—मॉडल कीसमम्य पूर्वमुमम तक पहुँच, उम्मेद्वह रक्छिड और attacker काअपना संदर्भ मॉडल—किसीhow-to के रूप में नहीं बल्कतकि threat पर अस्पष्ट डर कीबजय तर्क कियाजासके।

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.



The privacy signal is hidden in an ordinary prediction request.

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

हमलाकोकिसीवशिष गोपनीयताAPI कीजरूरत नहीं। यदमॉडल पहले देखे गए रक्छिड पर असमम्य वश्वस दखिताहै तोसमम्य नैदमकि प्रश्न सदस्यतासंकेत रसिम कर सकतीहै।

Original Aurora diagram — The Clean Paper CC BY 4.0

उन्हें क्यामला

तमि नष्किर्ष, और हर अगलीपहले से अधिक तीखी।

Averages उजगर लोगोकोछपादेते हैं। समग्र में—यमीआम तरीका—इनमें से कई मॉडल आश्वस्त रूप से private देखते हैं: अधकिंश मरीज के लिए हमलासंयोग से बेहतर नहींकरता। Per-patient view ने अलग कहनीदखिई। Knolle ने अध्ययन announcement में कहाकिपहले के assessments केवल सभीमरीज काऔसत जोखिम मपते थे; पहलीबार व्यक्तिमरीज के स्तर पर देखने से तस्वीर बहुत अलग नकिली। कुछ लोगोके लिए हमलालगभग पूरीतरह सफल होताहै—लेखक इसे 0.95 याअधिक हमलाAUC के रूप में मपते हैं (0.5 coin toss, 1.0 पूर्ण)—जबकि पूरे डेटसेट काऔसत संयोग से बेहतर नहीं देखता।



औसत संयोग के पस रह सकताहै, जबकि रिकॉर्ड कीछोटीtail बहुत अधिक उजगर रहतीहै। शोधपत्र काबदु है कि गोपनीयतामरीज level पर जाँच करनीहोगी केवल समग्र में नहीं।

Original Aurora diagram — The Clean Paper CC BY 4.0

एक्सपोजर असमन है—और कम प्रतिनिधित्व वालालोगोपर ज्यादापड़ताहै। पस-perfect हमलाके प्रति सबसे vulnerable मरीज व्यवस्थिति रूप से उन समूह में थे जोडेटामें कम प्रतिनिधित्व वालेथे: minority ethnicities, rare-बीमरी phenotypes, unusual इमेजिंग characteristics। Electronic-record डेटसेट में Black मरीज सबसे vulnerable रिकॉर्ड के बीच अपेक्षति से 31% अधिक बार आए; mammography set में malignancy के लिए suspicious scans उस danger क्षेत्र में आश्चर्यजनक +1,179% overrepresented थे। (ये दोचतिर उदहरण हैं, पूरीतस्वीर नहीं—शोधपत्र कई समूह में यही skew देखताहै—और ये छोटीextreme-risk tail के भीतर समेक्ष over-निरूपण हैं: उजगर रिकॉर्ड में ये मरीज अपेक्षासे कतिनीअधिक बार आते हैं, न कि Black याsuspicious-mammography मरीज काकतिनाहसिंसाउजगर है।) मॉडल के पस इन मरीज जैसे कम उदहरण होते हैं जनिमें उन्हें “मिला” सके, इसलिए उनकारिकॉर्ड अलग दखिई देताहै—और अलग दखिनावहीचीज है जिसे सदस्यताहमलापकड़ताहै। गोपनीयताfailure यहृच्छकि नहींहै; वह margins पर पहले से मौजूद लोगोपर केंद्रति होतीहै।

बड़े मॉडल समस्याबढ़ते हैं। पस-perfect हमले के लिए उजगर मरीज कीसंख्यामॉडल क्षमताके सग तेजीसे बढ़ी—एक dermatology डेटसेट में हसिंसासबसे छोटे मॉडल में लगभग शून्य से बढ़कर सबसे बड़े में करीब दस में एक होगया। चकित्सकीय AI मॉडल जतिने बड़े और capable होते जा रहे हैं—पूरे क्षेत्र कीदशायहीहै—लेखक चेतमनीदेते हैं कि यह खस जोखिम कम नहीं अधिक गंभीर होताहै।

Rückert का सर सीमा है: यह tolerable जोखिम नहीं है; स्वस्थ डेटा अत्यंत sensitive है।

“औसतन सुरक्षित” गलत परीक्षण क्यों है

कम औसत जोखिम को सफ़ा bill of स्वस्थ समझने का आकर्षण है। शोधपत्र का मूल सबक है कि यहाँ averaging केवल imprecise नहीं—गलत चीज़ मान रही है।

गोपनीयता गारंटी भी सही है जब वह सबसे अधिक जोखिम वाले व्यक्तियों के लिए भी काम करे, केवल बीच वाले के लिए नहीं। ऐसामॉडल जिनमें 1,000 में 999 मरीज की सदस्यता recover न हो सके लेकिन एक व्यक्ति लगभग पूर्ण तरीके से identify हो सके, उस एक व्यक्ति के लिए किसी अर्थ में “99.9% private” नहीं है। और यदि वह एक व्यक्ति predictably rare-बैथनी मरीज या ethnic-minority मरीज है, तो मेट्रिक केवल अधूरा नहीं बल्कि चुपचाप discriminatory है। वह बहुमत की सुरक्षामता है और उसे सबकी सुरक्षा कह देता है।

यही बदलाव शोधपत्र मजबूर करता है: यह मॉडल औसत में कितना private है? से सबसे उज्जगर मरीज कैसे है, और वह किस समूह से है? तक। ये अलग सवाल हैं, और गोपनीयता के लिए दूसरा ही असली सवाल है।

यह क्या सद्धि या दानही करता

- यह नहीं कहता कि चकित्सकीय AI छोड़ देनी चाहिए। लेखकों को फ़रेमिंग mitigation है, retreat नहीं जोखिम मंजूर और ठीक करें, building बंद न करें।
- इसका अर्थ यह नहीं कि हर चकित्सकीय मॉडल रसिन्न कर रहा है या कोई खास तैनात मॉडल हमला हो चुका है। यह दिखाता है कि जोखिम मौजूद है और असमम रूप से बाँटा है, और मजबूत समग्र मेट्रिक्स उसे पूरान कर पता करते हैं।
- इसका अर्थ यह नहीं कि आपके रिकॉर्ड पहले से उज्जगर हैं। हमला जोखिम स्थितियाँ चाहिए—मॉडल पहुँच, उम्मीदवार रिकॉर्ड और attacker की अपनी infrastructure—यह casual क्षमता नहीं है।
- यह मरीज content रसिन्न होना नहीं दिखाता। रसिन्न सदस्यता है—शमलि होने का तथ्य—जो अलग वजह से खतरनाक है, न कि रिकॉर्ड बहर आ जाने के कारण।
- यह disparities को एक सुरक्षित संख्या में नहीं समेटता। मजबूत reproducible परणिस पैटर्न है—समग्र मेट्रिक्स व्यक्ति जोखिम को कम देखते हैं, और कम प्रतिनिधित्व वाला मरीज सबसे ज्यादा भाग उठाते हैं—जो डेटासेट और सैकड़ों मॉडल में दिखा।

सक्ष्य कतिना मजबूत है?

यह empirical, पद्धतगित परणिस है और मजबूत है। पैटर्न—समग्र गोपनीयता मेट्रिक्स per-patient जोखिम को systematically कम देखती हैं, बचा हुआ जोखिम कम प्रतिनिधित्व वाला समूह पर केंद्रित होता है, और मॉडल क्षमता के साथ बढ़ता है—सब अलग डेटा types वाले डेटासेट और हर डेटासेट में बड़ी संख्या में मॉडल पर काम रहा। यही व्यवस्था इसे एक-डेटासेट artefact की वजह से credible मान देना बनती है।

दोई मजबूत caveats हैं। पहला यह जोखिम का प्रदर्शन है, लेखक द्वारा बनाए हमले चलकर मारा गया यह बताता है कि stated मजबूत के तहत capable attacker क्या कर सकता है, न कि वास्तविक दुनिया हमले कतिनी बार होते हैं। दूसरा vivid

संख्याएँ—कुछ मरीज के लिए पस-perfect success—डजिइन के अनुसार सबसे worst-off व्यक्ति बतते हैं; यही पूराबिंदु है। उन्हें “tail औसत के संकेत से कहीं भी नहीं” की तरह पढ़ना चाहिए, “most मरीज उज्जर हैं” की तरह नहीं।

उचित रुख न चेतनी है, न dismissal: कई डेटासेट पर किया गया सख्त अध्ययन दिखाता है कि medical-AI गोपनीयता certify करने का मतलब कि आपने worst मामले को देख ही नहीं पता—और वही worst मामले उन मरीज पर पड़ते हैं जिनके पास पहले से सबसे कम margin है।

यह क्यों मानने रखता है

दोबतें इसे technical footnote से अधिक बनती हैं।

पहली यह बदलता है कि “privacy-preserving” कहने की अनुमति किसे मिलनी चाहिए। यदि मॉडल औसत गोपनीयता स्कोर के साथ रलीज किया जाता है, स्कोर सचमुच कम हो सकता है और फिर भी मरीज का subset सदस्यता हमला से पस-perfectly identifiable हो सकता है। शोधपत्र की वजह से कि demand ठोस है: रलीज से पहले गोपनीयता जोखिम हर मरीज के स्तर पर assess करें, तैनात मॉडल की पहुँच नियंत्रित करें, और अंतर गोपनीयता जैसी techniques इस्तेमाल करें—प्रशिक्षण में थोड़ा carefully calibrated शोर जोड़ना जो सदस्यता हमले को blunt करता है, मॉडल usefulness की वस्तुविक और managed लगत के साथ। गोपनीयता-utility trade-off स्पष्ट है, मुक्त नहीं। “हमने औसत जैसा किया अब पर्याप्त नहीं होना चाहिए।

दूसरी यह गोपनीयता को fairness से जोड़ता है। चकित्सकीय डेटा में जो समूह कम प्रतिनिधित्व वला हैं—और इसलिए चकित्सकीय AI से पहले ही खरब serve हो सकते हैं—वही समूह नकिलते हैं जिनकी गोपनीयता AI सबसे कम सुरक्षा देती है। क्षेत्र पहली inequity पर काम करते हुए दूसरी को किसी और वभिन्न कामुद्दान हीमान सकता लगे वही है।

Medical-AI गोपनीयता की आश्वस्त कहनी—हमने माँ किया औसत कम है, सब ठीक है—यही कहनी शोधपत्र खेलकर अलग करता है। प्रैक्टिकी से किसी को डरकर भगने के लिए नहीं बल्कि मिनक को वहलाने जमे के लिए जहाँ वह होना चाहिए: ऐसा promise जिसे आप तभी दवा कर सकते हैं जब आपने उसे उस व्यक्ति के लिए जैसा किया हो जिसके hurt हमें की संभनना सबसे अधिक है।

साफ सारांश

सदस्यता अनुमान हमले यह पता लगाने की कोशिश करते हैं कि किसी खास व्यक्ति का किण्ड AI मॉडल के प्रशिक्षण डेटा में था या नहीं—और क्योंकि सदस्यता खुद revealing हो सकती है (जैसे आपको कोई खास बीमारी थी), यह वस्तुविक गोपनीयता रसिध है भले मूल रकिण्ड कभीबहर न आए। पहले का काम डेटासेट पर ऐसे हमले की औसत सफलता मपता था। इस अध्ययन ने सत चकित्सकीय डेटासेट—इमेजिंग, ECG, इलेक्ट्रॉनिक रकिण्ड—और हर डेटासेट के कई मॉडल में इसे हर मरीज के लिए अलग मपता और पन्नाकिसमग्र मेट्रिक्स जोखिम को बहुत कम दिखाती हैं: कुछ व्यक्ति लगभग perfectly identify किए जा सकते हैं (हमला $AUC \geq 0.95$) भले डेटासेट-wide औसत सुरक्षित लगे; सबसे vulnerable systematically कम प्रतिनिधित्व वला समूह—minority ethnicity, दुर्लभ बीमारी unusual इमेजिंग—से हैं; और मॉडल आकार के साथ समस्या बढ़ती है। लेखक चकित्सकीय AI छोड़ने को नहीं कहते—वे individual-level गोपनीयता मानन, मॉडल पहुँच नियंत्रण और अंतर गोपनीयता की माँग करते हैं। Takeaway: “औसत में private” गोपनीयता गारंटी नहीं है, और जहाँ यह वफिल होना होती है वहाँ अक्सर वही मरीज होते हैं जिन्हें पहले से सबसे कम सुरक्षामिली है।

बनाबदाचदकर जाँच

पेपर क्यादखिताहै: सत चकित्सकीय डेटासेट और हर डेटासेट के कई मॉडल में कुछ व्यक्ति के लिए per-patient membership-अनुमति जोखिम समग्र मेट्रिक्स के संकेत से बहुत अधिक है—परफेक्ट हमला (AUC ≥ 0.95) तक—और बचाहुआ जोखिम कम प्रतिनिधित्व बलासमूह पर disproportionate रूप से पड़ताहै तथा मॉडल क्षमताके साथ बढ़ताहै।

क्यासंभव है लेकिन सिद्ध नहीं कि वास्तविक दुनिया attackers तैनात क्लिनिकल मॉडल के विरुद्ध इसे पहले से इस्तेमाल कर रहे हैं; अध्ययन stated मजबूती में क्षमता और उसका distribution देखतीहै, वास्तविक दुनिया में हमले की frequency नहीं।

यह क्या नहीं देखिता कि चकित्सकीय AI छोड़ देनी चाहिए; कि हर मॉडल रसिप्त करताहै या कोई विशिष्ट तैनात मॉडल breach हो चुकाहै; कि patient-record contents सदस्यता की बजाय उजागर करनाहोते हैं; या disparity की एक ही quantified संख्या—मजबूत परिणाम पैटर्न है, कोई एक संख्या नहीं।

मुख्य सीमाएँ: लेखकों के अपने हमले से worst-case जोखिम मिला गया देखा गया incidents नहीं dramatic चित्र डिज़ाइन के अनुसार सबसे उजागर व्यक्ति बतते हैं; और सटीक समूह disparities को एक संख्या में समेटने की बजाय डेटासेट में संगत पैटर्न के रूप में दिखाया गयाहै।

समग्र पठक को कतिना भरोसा रखना चाहिए? इस पर उच्च भरोसा कि समग्र गोपनीयता मेट्रिक्स व्यक्ति जोखिम को कम देखतीहै, बचाहुआ जोखिम कम प्रतिनिधित्व बलासमूह पर केंद्रित होताहै और मॉडल आकार के साथ बढ़ताहै—यह कई डेटासेट और मॉडल में दिखाया गयाहै। ऐसे हमले वास्तविक दुनिया में कतिनीबार होते हैं, इस पर मध्यम भरोसा क्योंकि अध्ययन इसे मजबूती नहीं सुरक्षित पढ़ें: न “चकित्सकीय AI आपका डेटा रसिप्त करतीहै,” न “गोपनीयता हल करना हो गई”; बल्कि “मजबूत गोपनीयता जाँच अपने worst मामले को नहीं देखती और वे मामले सबसे vulnerable मरीज पर पड़ते हैं—इसलिए जाँच बदलनी होगी”

स्रोत

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>