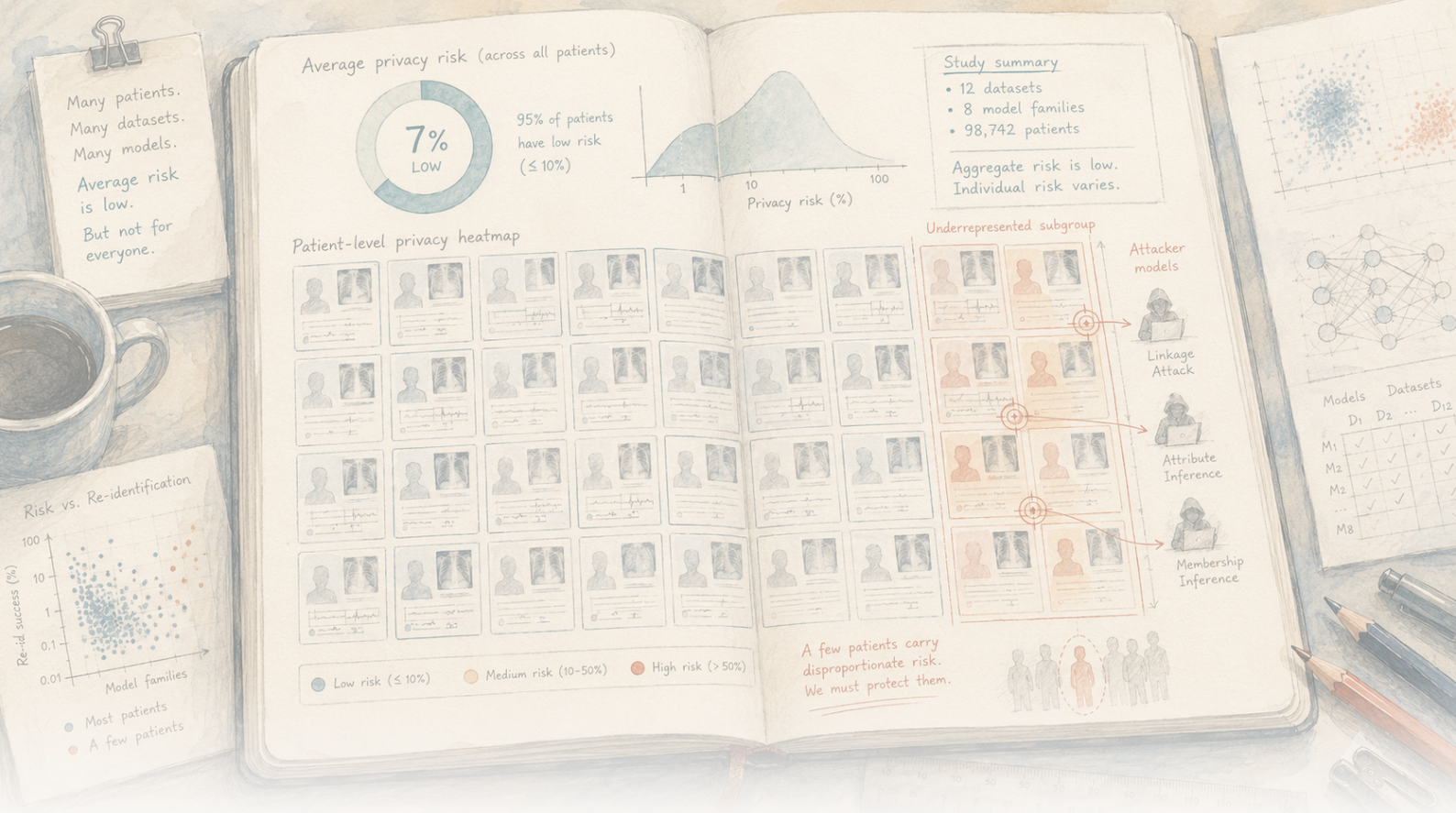




The CLEAN PAPER

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

Medical AI average-এ private হয়েও নরিদষ্টি patient-কে expose করতে পারে — বিশেষ করে underrepresented patient-কে

Seven medical dataset ও বহু model-এ per-patient membership-inference analysis দেখায় aggregate privacy metric severe individual risk hide করতে পারে। কিছু patient-এর attack AUC ≥ 0.95 হলেও overall average chance-level, এবং extreme-risk tail-এ minority ethnicity, rare phenotype ও unusual image disproportionately বেশি। Model capacity বাড়লে tail risk বাড়ে। Study patient record content leak বা actual deployed attack frequency measure করে না; এটি দেখায় “private on average” ব্যক্তি-level guarantee নয়। Authors per-patient testing, access control ও differential privacy-এর মতো mitigation চান।

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

গণ্যপনীয়তা নশ্চয়তা গড়কে নয়, ব্যক্তির রোগীকে দেওয়া প্রতিশ্রুতি

চিকিৎসাবিষয়ক AI মডেলকে “গণ্যপনীয়তা-সংরক্ষণকারী” বলা হলে প্রায়ই একটি সামগ্রিক গড় দেখানো হয়: প্রশিক্ষণ-তথ্যে থাকা রোগীদের মধ্যে গড় সদস্যত্ব-অনুমতি সন্ধিধান্ত ঝুঁকি কম। শুনতে আশ্বস্তকর। এই গবেষণাপত্রের অস্বস্তিকর দিক: **গড় ভুল প্রশ্ন হতে পারে।**

গণ্যপনীয়তা একটি জনসংখ্যা গড় নয়; প্রতিটি ব্যক্তির জন্য প্রতিশ্রুতি—প্রশিক্ষণ সটে ছিল বলে যেন পরে তাকে শনাক্ত করা না যায়। গড় 999 রোগীকে সুরক্ষিত করে। জনকে প্রায় সম্পূর্ণভাবে উন্মুক্ত করেও নমিন দেখাতে পারে। গবেষকরা প্রথমবার প্রতিটি রোগীকে আলাদা করে ঝুঁকি পরিমাপ করে দেখেছেন: সামগ্রিক হিসাবে নিরাপদ বলে মনে হওয়া মডেল নির্দিষ্ট ব্যক্তির সদস্যত্ব প্রায় নিখুঁতভাবে ফাঁস করতে পারে, এবং উন্মুক্ত ব্যক্তি অসমভাবে কম প্রতিনিধিত্বশীল গণ্যপনীয় সদস্য।

গবেষকরা কী করছেন

Moritz Knolle, Daniel Rückert, Georg Kaissis ও সহকর্মীরা ভালোভাবে-প্রচিতি সদস্যত্ব অনুমতি সন্ধিধান্ত আক্রমণ-এর প্রশ্ন বদলান। “ডটোসটে গড়ে আক্রমণ কতবার সফল?”-এর বদলে “**এই রোগী কত উন্মুক্ত?**”।

চিকিৎসাবিষয়ক ইমজিং, ECG ও ইলেকট্রনিক স্বাস্থ্য নথি—ভিন্ন তথ্য ধরনের 7 প্রতিষ্ঠিত চিকিৎসাবিষয়ক ডটোসটে-এ, প্রতিডটোসটে বহু মডেল (গবেষণাপত্রের বনিয়াসে 200 মডেল পর্যন্ত), রোগী-by-রোগী সদস্যত্ব ঝুঁকি অনুমান করা হয়। পরে রোগ অবস্থা, স্ব-প্রতিবেদিত জাতগত পরিচয়, বমি, লুগি, ইমজিং প্রোটোকল ইত্যাদি দল অনুযায়ী ভাগভিত্তিক বিশ্লেষণ।

সদস্যত্ব অনুমতি সন্ধিধান্ত আক্রমণ আসলে কী

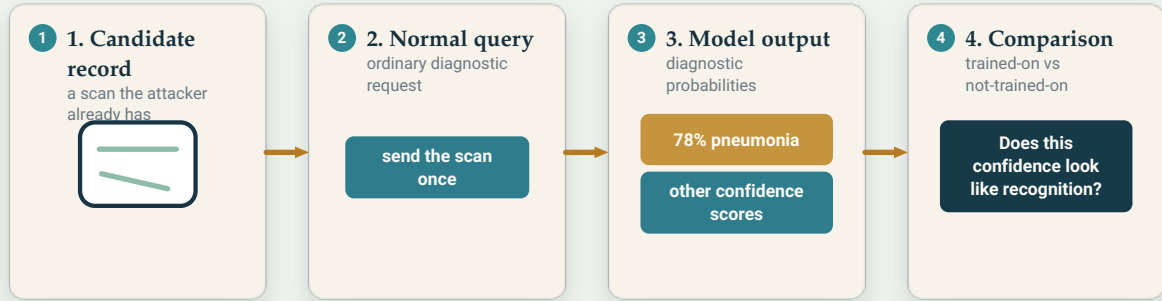
এখানে ফাঁস করা মানে মডেল “আপনার চিকিৎসাবিষয়ক নথি print করে দেয়” নয়। ফাঁস করা হলে **সদস্যত্ব**—আপনার নথি প্রশিক্ষণ সটে ছিল কিনা।

একটি রোগ নির্ণয়মূলক মডেলকে স্ক্যান দিলে সে নড়িমেঁটানি, cardiomegaly ইত্যাদির সম্ভাবনা দেয়। আক্রমণটি এই সাধারণ আউটপুটই ব্যবহার করে। প্রশিক্ষণে দেখা উদাহরণে মডেলটি অদেখা উদাহরণে তুলনায় সামান্য বেশি আত্মবিশ্বাসী হতে পারে—যেমন পরীক্ষার প্রশ্ন আগে দেখা কখনো শিক্ষার্থী। আক্রমণকারীর কাছে সম্ভাব্য রোগীর নথি থাকে; সে মডেলকে স্বাভাবিক রোগ নির্ণয়মূলক প্রশ্ন করে এবং আস্থার ধরনটি প্রশিক্ষণ-সদস্যের মতো কিনা দেখে। রফোরেন্স বা shadow model ব্যবহার করে “আগে দেখা” বনাম “অদেখা” উদাহরণে আস্থা-আচরণ শেখা যায়।

Request দেখতে স্বাভাবিক ক্লিনিক্যাল query-এর মতো; সংকেত পূর্বাভাস আস্থার ভেতর। কনে সদস্যত্ব সংবেদনশীল? যদি মডেল বিশেষ ক্যানসার চিকিৎসা দেওয়া রোগীর তথ্যে প্রশিক্ষণ দেওয়া হয়, কারণে সদস্যত্ব অনুমান করা তার রোগ/চিকিৎসা সম্পর্কে সংবেদনশীল তথ্য reveal করতে পারে—এমনকি যদি নথি বিষয়বস্তু বের না হয়। গবেষণা অনুমানগুলো—মডেল পূর্বাভাস প্রবশোধকির, প্রারথী নথি, আক্রমণকারী রফোরেন্স অবকাঠামো—থলে লাখুলি দিয়ে ঝুঁকি বিবেচনামূলক চিন্তার জন্য।

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.

Special গণনীয়তা API দরকার নই: সাধারণ রোগ নির্ণয়মূলক পূর্বাভাসে প্রশিক্ষণ উদাহরণে ওপর অস্বাভাবিক আস্থা সদস্যত্ব সংকতে তৈরি করতে পারে।

Original Aurora diagram — The Clean Paper CC BY 4.0

তারা কী পয়েছেন

গড় মান ঝুঁকরি লজে আড়াল করতে পারে। সামগ্রিকি মটেরকি বহু model-এর বহুদধে membership attack অধিকাংশ রোগীর জন্য দৈব অনুমানের কাছাকাছি। কিন্তু ব্যক্তি ধরে দেখলে কছু রোগীর ক্ষেত্রে attack AUC প্রায় নখিত— ≥ 0.95 ; এখানে 0.5 মান coin flip, 1.0 মান নখিত বিভাজন। অর্থাৎ dataset-wide গড় নরীহ দেখলেও নির্দিষ্ট কয়েকজনরে ঝুঁকি খুব বেশি হতে পারে।

A safe average can hide the exposed tail

Most patients cluster near chance. The privacy question lives in the small tail of records an attack can identify much more confidently.



গড় দবৈ সম্ভাবনার কাছে থাকতে পারে, অথচ ছোট subset-এর নথি অনেক বেশি উন্মুক্ত। গোপনীয়তা রোগী-স্তরে যাচাই না করলে লেজে অদৃশ্য থাকে।

Original Aurora diagram — The Clean Paper CC BY 4.0

ঝুঁকি অসম, কম প্রতিনিধিত্বশীল গোষ্ঠীতে কেন্দ্রীভূত করা। চরম-ঝুঁকি রোগী ধারাবাহিকভাবে সংখ্যালঘু জাতগিত পরিচয়, বরিলরোগ ফনেটাইপ, অস্বাভাবিক ইমজিং বৈশিষ্ট্যের মতো কম প্রতিনিধিত্বশীল গোষ্ঠীতে বেশি। ইলেকট্রনিক-নথি ডটোসটে কৃষ্ণ রোগী অধিকাংশ-ঝুঁকিপূর্ণ লেজে প্রত্যাশার তুলনায় 31% বেশি দেখা যায়; ম্যামোগ্রাফি ডটোসটে malignancy-suspicious স্ক্যান ঝুঁকি লেজে +1,179% আত্মীয় অতিরিক্ত প্রতিনিধিত্ব। এগুলো ওই দলের কত শতাংশ উন্মুক্ত—তা নয়; ছোট চরম-ঝুঁকি লেজে প্রত্যাশা গণনার তুলনায় অতিরিক্ত প্রতিনিধিত্ব।

প্রক্রিয়াগত অন্তর্দৃষ্টি: মডলের কাছে সদৃশ উদাহরণ কম হলে অস্বাভাবিক নথি বেশি “stand বাইরে” করতে পারে; সদস্যত্ব আক্রমণ সেই মুখস্থকরণ/overconfidence সংকতে শনাক্ত করে। গোপনীয়তা ব্যর্থতা এলে মিলে না।

বড় মডলে ঝুঁকি বাড়ে। মডলের সক্ষমতা বাড়ার সঙ্গে প্রায় নথিভাবে শনাক্তযোগ্য রোগীর সংখ্যা দ্রুত বাড়ে। একটি dermatology ডটোসটে সবচেয়ে ছোট মডলে উচ্চ-ঝুঁকি রোগী প্রায় ছিল না, কিন্তু সবচেয়ে বড় মডলে প্রায় প্রতি 10 জনে 1 জন উচ্চ-ঝুঁকি লেজে পড়ে। চিকিৎসাব্যয়ক AI কর্মে বড় হওয়ার প্রবণতায় এই নরিদ্রিষ্ট গোপনীয়তা-ঝুঁকিকে লেখকরা বাড়তে থাকা উদ্বেগে হিসেবে দেখেছেন।

কেন “গড়ে নরিপদ” ভুল পরীক্ষা

কম গড় ঝুঁকি মানই সবার গোপনীয়তা নরিপদ নয়। একটি গোপনীয়তা-নিশ্চয়তা শুধু “মধ্যম” রোগীর জন্য নয়, সবচেয়ে বেশি উন্মুক্ত রোগীর ক্ষেত্রেও অর্থবহ হওয়া দরকার। 1,000 নথি মধ্যে 999টি আলাদা করা না গেলেও যদি একটি নথিকে প্রায় নিশ্চিতভাবে শনাক্ত করা যায়, তবে “99.9% private” বলা সেই ব্যক্তির জন্য অর্থহীন। আর যদি উচ্চ-ঝুঁকি এই লেজে বরিল রোগ বা কম প্রতিনিধিত্বকারী গোষ্ঠীর রোগীরা বেশি থাকে, তবে গড়ভিত্তিক মটেরিকি শুধু অসম্পূর্ণ নয়—বৈষম্যও আড়াল করতে পারে।

গবেষণাপত্র উপস্থাপনভঙ্গিবিদলায়: মডলে গড়ে কত ব্যক্তিগত? থেকে সবচেয়ে উন্মুক্ত রোগী কে, এবং তারা কোন দলের?। দ্বিতীয় প্রশ্ন

গোপনীয়তা নশ্চয়তার সঙ্গে বেশি সামঞ্জস্যপূর্ণ।

এই গবেষণা কী প্রমাণ করে না

- চিকিৎসাবিষয়ক AI ভাগ করা করতে হবে—লেখকরা প্রশমন চান, retreat নয়।
- প্রতিটি মতে তায়নে করা চিকিৎসাবিষয়ক মডলে ফাঁস করছে বা breached হয়েছে—না; ঝুঁকিসিদ্ধতা দেখানো হয়েছে।
- আপনার নথি ইতিমধ্যেই উন্মুক্ত—না; আক্রমণের নরিদৃষ্টি প্রবশোধকার/নথি/রফোরনেস অনুমান আছে।
- রোগী বিষয়বস্তু বের হয়ে আসে—এটি সদস্যত্ব ফাঁস করা, নথি dump নয়।
- বৈষম্য একটি একক সর্বজনীন সংখ্যা—না; দুট ফল হলো multi-ডটোসটে ধরন।

প্রমাণ কতটা শক্ত?

প্রায়োগিক পদ্ধতিগত ফল হিসেবে শক্তিশালী। সামগ্রিক গড়ভিত্তিকি মটেরিকি ব্যক্তি ঝুঁকি understate করে, অবশিষ্ট ঝুঁকি কম প্রতিনিধিত্বশীল গোষ্ঠীতে কেন্দ্রীভূত করে এবং মডলে সক্ষমতার সঙ্গে worsens—ধরন 7 ভিন্ন চিকিৎসাবিষয়ক ডটোসটে ও বহু মডলে পুনরাবৃত্তি হয়েছে। Breadth এক-ডটোসটে ক্তরমি প্রভাব ব্যাখ্যা দুর্বল করে।

দুটি সতরকতা। প্রথমত এটি ঝুঁকি প্রদর্শন, বাস্তব জগতের ঘটনা কমপাঙ্ক নয়; লেখকরা নজিদের আক্রমণ উল্লিখিত অনুমানগুলোর অধীনে চালিয়ে সক্ষম আক্রমণকারী কতটা সফল হতে পারে দেখেছেন। দ্বিতীয়ত চোখে পড়ার মতো প্রায় নথিত সংখ্যা ইচ্ছাকৃতভাবে সবচেয়ে খারাপ-উন্মুক্ত লজে বর্ণনা করে; “অধিকাংশ রোগী উন্মুক্ত” নয়—বরং গড় লজে কতটা আড়াল করে সেটাই বিন্দু।

সঠিক অবস্থান alarmism বা খারজি নয়; মানক সামগ্রিক হিসাব গোপনীয়তা যাচাই সবচেয়ে খারাপ ক্ষেত্রগুলোর বাদ পড়া করে, এবং সবচেয়ে খারাপ ক্ষেত্রগুলোর সর্বনমিন-উপস্থাপতি রোগীতে বেশি পড়ে।

কেন এটি গুরুত্বপূর্ণ

প্রথম শিক্ষা: “গোপনীয়তা-সংরক্ষণকারী” বলার মানক বদলাতে হবে। প্রকাশের আগে প্রতিলিপি ঝুঁকি মূল্যায়ন করা, মডলে প্রবশোধকার ন্যিন্তরণ এবং কঠোর যমেন পার্থক্যমূলক গোপনীয়তা ববিচেনা করা দরকার। পার্থক্যমূলক গোপনীয়তা প্রশিক্ষণে ক্যালব্রিটেডে শব্দদূষণ/অ্যালগরিদমিকি সীমাবদ্ধতা দিয়ে সদস্যত্ব সংকতে কমায়, কিন্তু গোপনীয়তা-utility সুবধি-অসুবধিার বনিমিয় বাস্তব; মুক্ত সুরক্ষা নয়। “গড় যাচাই করা” যথেষ্ট নয়।

দ্বিতীয় শিক্ষা হলো ন্যায্যতার প্রশ্নের সঙ্গে গোপনীয়তাকেও একসঙ্গে দেখা দরকার। চিকিৎসা-তথ্যে কম প্রতিনিধিত্বশীল গোষ্ঠী প্রবাসসরে মানকে ক্ষতিগ্রস্ত হতে পারে; একই গোষ্ঠী যদি গোপনীয়তার দিক থেকেও সবচেয়ে কম সুরক্ষিত হয়, তবে দুটি বৈষম্যকে আলাদা বিভাগের সমস্যা হিসেবে দেখা যায় না।

সংক্ষিপ্ত সারাংশ

Membership-inference attack নরিধারণ করতে চায় কোনো নরিদর্শিট রোগীর নর্থি model training data-তে ছিল কিনা। সাতটি medical dataset—imaging, ECG ও EHR—এ ব্যক্তি-স্তররে ঝুঁকি মপে গবষণেটি দখোয়, সামগ্রিকি গড় ঝুঁকি বিভিন্নতকির হতে পারে: dataset-এর গড় দবৈ অনুমানরে কাছাকাছি হলেও কিছু ব্যক্তির attack AUC ≥ 0.95 । উচ্চ-ঝুঁকির এই লজে কম প্রতিনিধিত্বকারী গেষ্টী—সংখ্যালঘু জাতগিত পরচিয়, বরিল রোগ বা অস্বাভাবিকি imaging—এ বেশিকেন্দ্রীভূত; model বড় ও সক্ষম হলে সমস্যা আরও বাড়তে পারে। লখেকরো medical AI বন্ধ করতে বলনে না; তারা ব্যক্তি-স্তররে privacy audit, access control ও differential privacy ব্যবহাররে পক্ষে যুক্তদিনে। “গড়ে ব্যক্তিগিত” হওয়া কোনো ব্যক্তির গেপনীয়তার নশ্চয়তা নয়।

সরাসরি যাচাই

পপোর যা দখোয়: Multi-ডটোসটে/মডলে গবষণেয় প্রতরিগেী সদস্যত্ব ঝুঁকি সামগ্রিকি হিসাব গড়রে চয়েে কিছু ব্যক্তিতে অনেক বেশি; কম প্রতিনিধিত্বকারী দল অসমভাবে উন্মুক্ত; বড়তর মডলে ঝুঁকি বাড়ি।

যা সম্ভব, কনি্তু প্রমাণতি নয়: বাস্তব আক্রমণকারী মোতায়নে করা মডলে বর্তমানে একই আক্রমণ করছে—কম্পাঙ্ক গবষণে প্রমাপ করনে।

যা দখোয় না: প্রতটি চকিিসাবষিয়ক মডলে leaks, নরিদর্শিট breach, নর্থি বিষয়বস্তু প্রযবক্ষেণকাল, চকিিসাবষিয়ক AI ত্যাগ করা উচতি।

প্রধান সীমাবদ্ধতা: লখেকরো-তরৈ আক্রমণ উল্লখিতি অনুমানগুলোর অধীনে; সবচয়েে খারাপ-লজে মটেরিকি ইচ্ছাকৃতভাবে চরম ব্যক্তি মনোযোগ ধরে রাখা; সুনরিদর্শিট দল বৈষম্য ডটোসটে-নরিদর্শিট।

সাধারণ পাঠকরে আস্থা: সামগ্রিকি গড়ততিকি মটেরিকি ব্যক্তি ঝুঁকি বাদ পড়া করে এবং কম প্রতিনিধিত্বকারী রোগীর ওপর লজে কেন্দ্রীভূত করতে পারে—উচ্চ। বাস্তব জগতরে আক্রমণ কম্পাঙ্ক—মাঝারি/অজানা।

উৎস

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>