

A Federated, AI-Curated Clinical Knowledge Archive for Diagnosis and Research

Goal

Individual hospitals and clinicians generate a continuous stream of disease, diagnosis and treatment data, but no single institution can turn it into a systematic, reusable body of knowledge. This proposal outlines an AI-assisted pipeline that converts distributed, heterogeneous clinical records into a unified, freely accessible archive, which can serve both clinical diagnosis and biomedical research, and can act as empirical grounding for broader collective efforts to build comprehensive models of the human body.

Reducing the burden on clinicians

The central design choice is that AI, not medical staff, should carry the cost of standardizing data. Records arrive in incompatible formats, and a large fraction exist only on paper; requiring clinicians to re-encode them into a fixed schema is impractical and has repeatedly failed in past standardization efforts. Instead, a staged, self-bootstrapping training pipeline is proposed: (1) already-digitized records are used first to train baseline extraction models and build robust statistics; (2) non-digitized (paper) records are processed through OCR and related digitization methods; (3) the models trained in stage 1 are then used to complete and correct errors introduced in stage 2, so that noisy or partial digitized records are refined using patterns learned from clean data. This can be rolled out incrementally by medical department or specialty, allowing early, contained deployments. One caveat to track explicitly: institutions with fully digitized records are not a random sample of clinical practice, so biases in the stage-1 baseline could propagate when correcting stage-2 data; this should be monitored, not assumed away.

Governance and access

Processing is coordinated by a small number of independent centers rather than one central authority, reducing both single-point control and the risk of a single institution's biases dominating the archive. The resulting archive is freely accessible, and each entry is linked to relevant ongoing research through automated matching (e.g., embedding-based similarity) between archive content and the literature, keeping the archive connected to active science rather than acting as a static repository.

Privacy and quality control (collective-intelligence design)

Patient privacy is addressed through federated learning: raw records never leave the originating institution; only model updates or aggregated statistics are shared with the processing centers, and de-identification is applied locally before any aggregation. Quality control is designed to minimize added clinician workload while still generating a genuine incentive to participate: clinicians who draw on the archive for diagnosis or research give a simple score to the specific entries they used. This score serves two purposes at once — it flags unreliable entries for review by the specialized verification teams, and it routes recognition back to the institution or individual that originally contributed the underlying case, giving clinicians a concrete, low-effort form of credit for sharing data. Legal protections against misuse of the archive (commercial exploitation, copyright claims over aggregated clinical knowledge) are a necessary companion to this design; this is flagged here as a distinct problem best addressed with legal and policy scholars rather than resolved technically.

Why this needs a collective effort

No single hospital, lab, or country holds enough case diversity to detect rare disease patterns, diagnostic error trends, or treatment-response variation on its own. An AI-mediated, privacy-preserving, incentive-aware pipeline is what allows this fragmented knowledge to be pooled into a shared resource — turning routine clinical practice into a continuously growing, verifiable layer of evidence for both patient care and the collective modeling of human disease.