

Electronic Medical Records for Genetic Research: Results of the eMERGE Consortium

Abel N. Kho,^{1,2*} Jennifer A. Pacheco,¹ Peggy L. Peissig,³ Luke Rasmussen,³ Katherine M. Newton,^{4,5} Noah Weston,⁴ Paul K. Crane,⁶ Jyotishman Pathak,⁷ Christopher G. Chute,⁷ Suzette J. Bielinski,⁷ Iftikhar J. Kullo,⁸ Rongling Li,⁹ Teri A. Manolio,⁹ Rex L. Chisholm,¹ Joshua C. Denny¹⁰

Clinical data in electronic medical records (EMRs) are a potential source of longitudinal clinical data for research. The Electronic Medical Records and Genomics Network (eMERGE) investigates whether data captured through routine clinical care using EMRs can identify disease phenotypes with sufficient positive and negative predictive values for use in genome-wide association studies (GWAS). Using data from five different sets of EMRs, we have identified five disease phenotypes with positive predictive values of 73 to 98% and negative predictive values of 98 to 100%. Most EMRs captured key information (diagnoses, medications, laboratory tests) used to define phenotypes in a structured format. We identified natural language processing as an important tool to improve case identification rates. Efforts and incentives to increase the implementation of interoperable EMRs will markedly improve the availability of clinical data for genomics research.

INTRODUCTION

Electronic medical records (EMRs) have been promoted as essential to improving healthcare quality (1–4). Although current adoption rates remain low, recent government efforts may markedly increase the use of EMRs in clinical settings (5–9). The U.S. Centers for Medicare and Medicaid Services recently finalized a definition for “meaningful use” of EMRs, which defines standards for the recording and use of data in EMRs to promote quality care (10, 11). This standard, coupled with significant financial incentives and penalties, is intended to promote widespread adoption of EMRs within the U.S. healthcare system.

Understanding the strengths and limitations of current EMR data capture is crucial for identifying linkages between disease susceptibility and clinical presentation. In clinical care, EMRs serve to document clinical observations and patient-provider interactions and generate billing documentation. Clinical data captured in EMRs may have a secondary application in the research setting. In parallel with increasing EMR adoption, high-throughput DNA sequencing has made available millions of DNA sequence reads for genetic investigations (12). Understanding the current feasibility of linking clinical data captured in EMRs and genome sequencing data has important implications for genetics research and the promise of personalized medicine (13–15).

Genome-wide association studies (GWAS) require accurate classification of disease phenotypes to maintain adequate statistical power (16, 17). The Electronic Medical Records and Genomics Network (eMERGE) (18) aims to determine whether data captured through rou-

tine clinical care using EMRs can identify disease phenotypes with sufficient positive and negative predictive values (PPVs and NPVs) for application in GWAS. If successful, identification of disease phenotypes using EMR data may enable efficient and rapidly scalable genetic research. Specifically, greater efficiency may be gained by undertaking genome-wide single-nucleotide polymorphism (SNP) genotyping only once. EMR data may be used to determine whether the individual associated with each DNA sample is used as a case, a control, or neither for multiple phenotypes, facilitating a GWAS for each phenotype. The marginal cost of each GWAS after the initial genotyping expense is then limited to the costs involved with establishing and validating the operational EMR-based phenotype definition and the costs of performing the association analyses. Indeed, as genotyping costs continue to rapidly decline, efficient and cost-effective means to identify phenotypic data from EMRs takes on increasing importance (19).

However, it is unclear whether current EMR implementation captures clinical data adequately to identify patients for research aimed at identifying the genetic basis of disease susceptibility. The eMERGE consortium has a unique opportunity to evaluate the utility of current EMRs for genomic research and to identify key areas for improvement. Here, we determine whether data recorded in EMRs for routine clinical care at five U.S. study sites can be used to define phenotypes for genomic research, and discuss the challenges and lessons learned in using data extracted from existing EMRs for GWAS.

RESULTS

We analyzed EMR data collected from five eMERGE study sites to identify cases with one of five different disease phenotypes: dementia, cataracts, peripheral arterial disease, type 2 diabetes, and cardiac conduction defects. Table 1 lists the primary phenotypes, biorepository description, and EMR characteristics for each study site. Three sites used an internally developed EMR system for both inpatient and outpatient care; the remaining two sites used commercial EMR systems. One site used different EMR systems for inpatient and outpatient care.

¹Feinberg School of Medicine, Northwestern University, Chicago, IL 60611, USA. ²Regenstrief Institute Inc., Indianapolis, IN 46202, USA. ³Marshfield Clinic Research Foundation, Marshfield, WI 54449, USA. ⁴Group Health Research Institute, Seattle, WA 98101, USA. ⁵School of Public Health, University of Washington, Seattle, WA 98104, USA. ⁶School of Medicine, University of Washington, Seattle, WA 98104, USA. ⁷Department of Health Sciences Research, Mayo Clinic, Rochester, MN 55905, USA. ⁸Department of Internal Medicine, Mayo Clinic, Rochester, MN 55905, USA. ⁹Office of Population Genomics, National Human Genome Research Institute, Bethesda, MD 20892–2152, USA. ¹⁰Departments of Biomedical Informatics and Medicine, Vanderbilt University, Nashville, TN 37232, USA.

*To whom correspondence should be addressed. E-mail: a-kho@northwestern.edu

Some EMR systems captured data primarily from free-text documents (unstructured data), and others from a mix of structured data collection and free-text notes. Three sites used robust but different natural language processing (NLP) tools to extract structured data from free-text reports (20–25). Each study site had a separate DNA biorepository (linked to the EMR through a unique research identifier) to house biological samples for genotyping (26–29). With a single exception, all sites used an opt-in consent model to recruit participants into the biorepository. For our purposes, we analyzed only patients with records in both the institution’s biorepository and EMR.

Data required to define the clinical gold standard for the five selected disease phenotypes across study sites most commonly required only one category of data (for example, diabetes could be defined by laboratory tests alone, and peripheral arterial disease by a single radiological test), with one condition requiring two categories (Table 2). However,

algorithms to identify the same phenotypes using EMR data required multiple categories of data, ranging from one to four categories (for example, diagnostic information, medications, and laboratory tests), with additional data categories required to identify covariates and exclusion criteria. In the example of type 2 diabetes, the EMR-derived phenotype required diagnoses, laboratory tests, and medications to identify likely type 2 diabetes cases and used diagnoses to specifically exclude cases of type 1 diabetes. All sites used demographic, diagnoses, and medication data in their phenotype definitions.

The three study sites using internally developed, text-based EMRs required significant NLP efforts to extract concepts from free-text documents, with each using a different NLP platform. At these study sites, use of NLP tools enabled disease phenotype definitions using data stored in unstructured clinical notes (for example, ophthalmological examinations) and text-based reports [for example, radiology test results and

Table 1. Comparison of electronic medical records (EMRs) and biorepositories at five eMERGE institutions. GHC, Group Health Cooperative; NLP, natural language processing.

Institution	Biorepository overview	Recruitment model	Repository size (race/ethnicity)	EMR summary	Primary phenotype	Phenotyping methods*
Group Health (Seattle, WA)	GHC Biobank	Disease-specific cohort	4000 (>96% Caucasian)	Comprehensive vendor-based EMR since 2004	Dementia	Structured data extraction, free-text searches, manual chart review
	Alzheimer’s Disease Patient Registry and Adult Changes in Thought Study			20+ years pharmacy data 15+ years ICD9 data		
Marshfield Clinic Research Foundation (Marshfield, WI)	Personalized Medicine Research Project	Population-based	20,000 (98% Caucasian)	Comprehensive internally developed EMR since 1985	Cataracts	Structured data extraction, NLP, intelligent character recognition
	Geographically defined cohort within an integrated regional health care system			75% participants have 20+ years medical history		
Mayo Clinic (Rochester, MN)	Vascular Diseases Biorepository	Disease-specific cohort	3500 (>96% Caucasian)	Comprehensive internally developed EMR since 1995	Peripheral arterial disease (PAD)	Structured data extraction, NLP
	Mayo Clinic Non-Invasive Vascular Laboratory and Exercise Stress Testing Lab			40-year history of data extraction		
Northwestern University (Chicago, IL)	NUgene Project	Population-based	10,000 (12% AA, 8% Hispanic)	Comprehensive vendor-based inpatient (2001) and outpatient (1999) EMRs	Type 2 diabetes	Structured data extraction, free-text searches
	Northwestern-affiliated hospitals and outpatient clinics			20+ years ICD9 data		
Vanderbilt University (Nashville, TN)	BioVU	Population-based, opt-out consent model	92,000 (11% AA)	Comprehensive internally developed EMR since 2000	Cardiac conduction	Structured data extraction, NLP
	Vanderbilt Clinic, diverse outpatient population			35+ years medical history data		

*Structured data extraction refers to retrieving EMR data that have been stored in a predefined format.

electrocardiogram (ECG) reports]. Sites without NLP tools or experience with limited phenotype definitions included only data available in a structured format, and therefore readily extractable from the EMR.

Across all five study sites, the percent of data captured and stored in a structured format consistently met or exceeded the current “meaningful use” final rule requirements (that is, goals for structured data capture and use defined by the Office of the National Coordinator

to promote quality improvement using EMRs), with the notable exception of allergies and smoking status. Height, weight, and race/ethnicity, although satisfying the meaningful use requirements, demonstrated varying capture rates across sites. Only one institution with a vendor-based EMR had any data on family history stored in a structured format. At other sites, family history information was stored only in clinician notes and this information could not be extracted readily

Table 2. Data categories for defining a clinical gold standard, an EMR-derived phenotype, and covariates and exclusion criteria.

Primary phenotype	Data categories			Phenotyping methods
	Clinical gold standard	EMR-derived phenotype	Phenotype cohort (for example, covariates and exclusion criteria)	
Dementia	Demographics, clinical notes (clinician documentation of mental status and histopathological examination data)	Diagnoses, medications	Demographics, laboratory tests, radiology reports	Structured data extraction, free-text searches, manual chart review
Cataracts	Clinical notes (ophthalmologic examination)	Diagnoses, procedure codes	Demographics, medications	Structured data extraction, NLP, intelligent character recognition
Peripheral arterial disease	Radiology test results (ankle-brachial index or arteriography)	Diagnoses, procedure codes, medications, radiology test results	Demographics	Structured data extraction, NLP
Type 2 diabetes	Laboratory tests	Diagnoses, laboratory tests, medications	Demographics, laboratory tests, height, weight, family history, smoking history	Structured data extraction, free-text searches
Cardiac conduction	ECG measurements	ECG report results	Demographics, diagnoses, procedure codes, medications, laboratory tests	Structured data extraction, NLP

Table 3. Data completeness and type by common clinical categories. Meaningful use goal for % data recorded in EMR and type listed for comparison. S, structured data; U, unstructured data; M, mixture of structured and un-

structured data; GHC, Group Health Cooperative; MCRF, Marshfield Clinic Research Foundation; Mayo, Mayo Clinic; NU, Northwestern University; VU, Vanderbilt University.

	GHC	MCRF	Mayo	NU	VU	Meaningful use goal
Number	2790	19,771	3336	8161	81,952	
Age	100% (S)	100% (S)	100% (S)	100% (S)	100% (S)	>50% (S)
Gender	100% (S)	100% (S)	100% (S)	100% (S)	100% (S)	>50% (S)
Race/ethnicity	84% (S)	69% (S)	94% (S)	84% (S)	80% (S)	>50% (S)
Home address	100% (M)	100% (S)	96% (S)	100% (M)	N/A*	
Height	76% (S)	96% (S)	98% (M)	92% (S)	64% (M)	>50% (S)
Weight	79% (S)	91% (S)	98% (M)	96% (S)	78% (M)	>50% (S)
Blood pressure	59% (S)	97% (S)	62% (M)	97% (S)	80% (M)	>50% (S)
Diagnoses	100% (S)	100% (S)	85% (S)	99% (S)	91% (S)	>80% (S)
Laboratory tests	100% (S)	98% (S)	81% (S)	99% (S)	97% (S)	>40% (S)
Medication	100% (S)	99% (S)	95% (M)	97% (M)	87% (M)	>80% (S)
Allergies	54% (M)	39% (M)	50% (M)	94% (M)	83% (U)	>80% (S)
Smoking history (any)	86% (S)	77% (U)	94% (M)	73% (M)	>90% (U)	>50% (S)
Smoking history (detailed numeric)	54% (M)	N/A	94% (U)	12% (M)	0%	>50% (S)
Family history	20% (M)	(U)	(U)	36% (M)	(U)	

*Addresses are removed from the Vanderbilt biorepository in the de-identification process.

Table 4. Performance of algorithms to identify cases and controls from EMRs for five primary phenotypes. GHC, Group Health Cooperative; MCRF, Marshfield Clinic Research Foundation; Mayo, Mayo Clinic; NU, Northwestern University; VU, Vanderbilt University.

	GHC	MCRF	Mayo	NU	VU
Primary phenotype	Dementia	Cataract	Peripheral arterial disease	Type 2 diabetes	Cardiac conduction (quantitative trait)
EMR data sources to define phenotype	Diagnoses, medications	Diagnoses, procedures, medications	Procedure reports	Diagnoses, laboratory tests, medications	Diagnoses, laboratory tests, medications, ECG results
Method to validate EMR phenotype	Physician review*	Trained chart reviewers	Compared to clinical gold standard	Physician review	Physician review
Number of cases/controls	747/2043	2642/1322	1679/1657	756/777	2950
Biospecimen number	2790	19,771	3336	8161	81,952
% of total biospecimen pool	26.8	13.4	50.3	9.3	3.6
PPV (case/control) (%)	73	98/98	94/99	98/100	97

*Review team included two physicians, a psychometrician, a neuropsychologist, and a study nurse.

even with NLP. To define study phenotypes, no site required data categories with low rates of capture in the EMRs (allergies, family history). Despite variations in categories and completeness of data capture across sites (Table 3), four of the five study sites achieved PPVs of close to 100% for use of EMR data alone to identify their primary disease phenotype (Table 4). One site achieved a lower PPV of 73% using EMR data to identify cases with dementia. Absolute numbers of cases identified by EMR ranged from 747 to 2950 cases. Of sites with unselected noncohort biorepositories, rates of case identification ranged from 3.6 to 13.4% of the total eligible population. Sites using disease-specific biorepositories had case identification rates of 26.8 to 50.3% of the total population, which, after excluding known controls, represented 71 and 90% of the cases identified through prospective cohort collection. NPVs ranged from 98 to 100% for the three sites generating control cases using electronic algorithms.

To assess the additional benefit of NLP, we performed a comparison of the number of cases identified using structured data alone compared with that using both structured data and NLP at one site (Vanderbilt University). At this site, the use of NLP tools identified 129% more cases of QRS duration (2950 versus 1288) than did the use of structured data and string matching only, while maintaining a PPV of 97%.

DISCUSSION

In our study, data captured in EMRs for routine clinical care proved adequate to define five disease phenotypes across five different study sites with robust PPV and NPV. Encouragingly, several recent reports (30–32) demonstrate that GWAS based on EMR-derived phenotypes successfully replicated identification of genetic sequences associated with increased disease risk. Although we could achieve high PPVs using case identification algorithms based on data captured through routine clinical care, we note some attrition in the number of cases identified by this approach compared with disease-focused prospective case identification. In our study, electronic algorithms identified 71 and 90% of the possible cases within two prospectively collected disease cohorts. Reduction in case identification rates may be compensated for by the efficiency and scalability of electronic algorithms across EMRs.

Across the five unique EMRs, diagnosis codes, medications, and laboratory tests were readily extracted to identify phenotypes for GWAS. Race/ethnicity, family history, exposure history (for example, smoking), and environmental exposures were documented less frequently across all EMRs and, where present, often were captured in free-text form (for example, clinicians notes) and without consistent or standard nomenclature. Capturing interpreted test results that are typically not recorded as structured data elements (for example, arterial Doppler and ECG data) and clinician diagnoses (such as found on a problem list) generally required NLP. As a result, significant informatics efforts were required to tailor algorithms to each institution’s EMR to accurately identify each phenotype.

Both “home-grown” and commercial EMRs demonstrated high PPV rates across the primary phenotypes. Given the far wider population using commercial EMRs in routine clinical care, this finding suggests potential for broad dissemination of our approach to identify cases and controls for genetic analyses to achieve well-powered studies, although the impact of differences among commercial EMR systems is unclear. Regardless of EMR type, study sites leveraged strengths in EMR data quality and site-specific data extraction methods to optimize phenotyping algorithms, often using data categories with a high proportion of structured data at sites without NLP capacity.

Historically, institutions with significant free-text documentation in their EMRs developed or adapted robust NLP tools to extract data for further analysis (20, 33, 34). NLP enabled sites to improve case finding by searching across a wider range of EMR data categories. The observation that NLP tools allowed identification of 129% more cases than were identified using purely structured data and string matching only emphasizes the value of information captured in free text and is consistent with previous studies (35–37). As a consortium, eMERGE identified use of NLP to extract data from text documents as a critical tool to improve data quality for phenotyping. Sites with NLP experience shared best practices with other consortium sites to develop NLP capacity at all sites. However, in our study, even sites without NLP tools successfully identified their primary phenotype, and one site successfully replicated previously identified genotype-phenotype associations for five diseases, including type 2 diabetes (31). Certain phenotype identification algorithms, such as those for type 2 diabetes, were im-

Downloaded from stm.sciencemag.org on September 6, 2011

plemented without use of sophisticated NLP; other algorithms, such as those for identifying cardiac conduction problems, were implemented with a combination of NLP and structured data extraction. This variation reflected institutional informatics capacity and a bias toward selection of phenotypes using data captured in structured formats at sites without NLP capacity. Sites without NLP capacity may be limited to identifying phenotypes using only data categories captured in structured fields. Approaches using only structured data could still achieve comparable PPVs, but would have lower case identification rates. However, efficient access to data across the entire spectrum of clinical EMRs can compensate for lower identification rates to identify adequate numbers for genetic studies.

Some data categories consistently reflected low rates of structured data capture (Table 3). The EMRs in this study used Office of Management and Budget categories for race/ethnicity (38). Here, low rates of documentation of race and ethnicity in the EMRs are consistent with previous studies of routine physician practice (39). However, lower rates of race and ethnicity documentation in EMRs may not significantly affect subsequent genetic studies. For genetic studies, ancestry estimates derived from genotype data are often used in primary association analyses rather than self-reported race/ethnicity, although the latter adds important sociocultural information independent of genetic ancestry that may be useful in more refined analyses (40). Similarly, in our study, family history was primarily documented in clinician notes and was not readily extracted even with NLP tools. One site with a vendor-based EMR featured a family history section enabling a mixture of structured and unstructured data capture, but attracted low rates of physician documentation. Our findings are consistent with previous studies, although current efforts are under way to promote standardized collection of key elements of family history within EMRs (41–43).

Environmental exposures play a significant role in expression of disease in genetically susceptible populations (44–47). Unfortunately, environmental factors, such as exposure to environmental toxins or contaminants, are rarely captured in existing EMRs, with the notable exception of smoking status. Substantial improvements in methods to collect and link environmental data to clinical data in EMRs may enable future studies of the association between disease and environment (48).

In our chart review, we identified a number of common data quality issues. Foremost, the absence of information may not reflect the absence of condition. Depending on the institution, significant care might be rendered at outside institutions and therefore would not appear in the study site's EMR. To address this limitation, we defined minimum data requirements (for example, two documented clinical visits) to enhance the opportunity for clinical documentation beyond a single visit. We encountered instances of structured results violating acceptable ranges of possibility (for example, a weight of 1000 kg and a height of 15 cm), requiring post-extraction censoring of impossible values. Lack of data equivalency posed challenges in merging data within a single EMR and across EMRs. Often, data are imprecisely labeled such that different measures might be inappropriately mixed together. For example, laboratory tests with similar names (for example, glucose) might represent different tests (for example, blood glucose concentration versus urine glucose concentration). Similarly, diagnostic certainty differed depending on whether the diagnoses were entered in clinical notes or for billing purposes and differed across sites due to local billing practices (49). We identified use of data standards for EMR documentation as a necessary foundation to improve data quality and achieve data equivalence across sites. As a consortium, we used

the federally endorsed Consolidated Health Informatics (CHI) standards (LOINC, ICD9/SNOMED, and RxNorm) to promote data equivalency, and facilitate data sharing between sites (50–52). Phenotyping algorithms most commonly included diagnosis codes, medications, and laboratory tests, which are well covered by the CHI standards ICD9, RxNorm, and LOINC, respectively.

Our study sites represented academic medical centers or institutions with significant research programs and may have a greater focus on rigorous data collection for potential future research, limiting the generalizability of our findings to non-research-oriented clinical care settings. However, recent national initiatives may promote more complete and standardized data collection across EMR-enabled clinical care settings. Greater adherence to standardized data collection may facilitate the role of EMRs in research and enable the sharing of phenotype definitions across EMR systems. The Centers for Medicare and Medicaid Services and the Office of the National Coordinator have written regulations defining meaningful use of EMRs that promote the recording of structured data and define coding standards for data categories such as diagnoses, laboratory tests, and medications. Clear documentation in EMRs is a necessary goal to achieve meaningful use and enables measurement and improvement in quality of care. Achieving this goal likewise improves the quality and volume of data available for research. Significant financial incentives for achieving meaningful use of an EMR (up to \$63,750 per provider over 4 years) may increase the future availability of structured and standardized data from EMRs. Although EMR data may not capture the nuance of the human-human interaction between patient and provider, accurate and structured capture of diagnosis, laboratory test, and medication data, supplemented with text mining tools, has proved useful for identifying disease phenotypes for GWAS within the eMERGE network.

Widespread adoption of EMRs creates the potential for a quantum shift forward in the availability of longitudinal, real-world clinical data for genetics research. Our study suggests that current EMRs used for routine clinical care can be used to identify phenotypes for genetic studies. Future investment in the dissemination, standardization, and comprehensive capture of phenotypic and environmental data in EMRs will help to achieve rapidly scalable phenotyping efforts to match the proliferation of genomics data.

MATERIALS AND METHODS

Using data from their EMR, each member of the eMERGE consortium selected a primary study phenotype and developed algorithms to identify the phenotype. We characterized EMRs as either internally or commercially developed and quantified the historical extent of data collection and primary methods and tools available to define phenotypes from the EMR (Table 1). We identified the primary consent model, recruitment numbers, and demographics of each biorepository. All sites received approval from their institutional review board for the conduct of this study.

We identified categories of EMR data used to define the five primary phenotypes (Table 2). At four of the five sites, as part of biorepository enrollment, additional data were collected on patients through an enrollment questionnaire (that is, additional data collection outside of the clinical EMR); the fifth site (Vanderbilt University) used an opt-out, de-identified collection model that precluded collection of biorepository-specific information.

For each data category, we generated a measure of data completeness, defined as percent of the cohort with at least one recorded entry within the EMR for each data category. We classified the type of data in each category as structured, unstructured (predominantly free text), or mixed. We defined structured data as numeric data or text data captured and stored in a predefined format as consistent with the current meaningful use definition. Unstructured data refer to data fields (for example, clinical notes) that typically require subsequent processing to be useful for phenotype identification algorithms. To identify a comparable cohort in each EMR, we defined study patients as those enrolled within the site's biorepository who had at least two in-person visits to the healthcare institution documented within the EMR. For the analyses presented here, study patients were not limited to those with one of the primary phenotypes.

To determine the accuracy of defining phenotypes using EMR data alone, we reviewed 100 clinical charts from the EMR at each site. Three sites used clinician chart review as the standard to confirm the primary phenotype from the records. One site used the clinical gold standard for their primary phenotype. The remaining site used trained EMR chart abstractors to confirm the primary phenotype. We measured the PPV of EMR data to correctly identify cases for the primary phenotype compared with chart review (the standard). For three of the five phenotypes, we measured the NPV of EMR data to correctly identify control cases for the primary phenotype compared with the chart review standard. One of the study sites measured a quantitative trait (QRS duration, a measure of cardiac conduction) precluding measurement of an NPV. For the remaining phenotype—dementia—sufficient research quality control subjects were available from an ongoing prospective cohort study, and there was concern that reliable identification of controls from EMR data would be prohibitively difficult (53, 54).

REFERENCES AND NOTES

- B. Chaudhry, J. Wang, S. Wu, M. Maglione, W. Mojica, E. Roth, S. C. Morton, P. G. Shekelle, Systematic review: Impact of health information technology on quality, efficiency, and costs of medical care. *Ann. Intern. Med.* **144**, 742–752 (2006).
- J. A. Linder, J. Ma, D. W. Bates, B. Middleton, R. S. Stafford, Electronic health record use and the quality of ambulatory care in the United States. *Arch. Intern. Med.* **167**, 1400–1405 (2007).
- M. N. Walsh, C. W. Yancy, N. M. Albert, A. B. Curtis, W. G. Stough, M. Gheorghiade, J. T. Heywood, M. L. McBride, M. R. Mehra, C. M. O'Connor, D. Reynolds, G. C. Fonarow, Electronic health records and quality of care for heart failure. *Am. Heart J.* **159**, 635–642.e1 (2010).
- R. J. Baron, Quality improvement with an electronic health record: Achievable, but not automatic. *Ann. Intern. Med.* **147**, 549–552 (2007).
- A. K. Jha, C. M. DesRoches, E. G. Campbell, K. Donelan, S. R. Rao, T. G. Ferris, A. Shields, S. Rosenbaum, D. Blumenthal, Use of electronic health records in U.S. hospitals. *N. Engl. J. Med.* **360**, 1628–1638 (2009).
- C. M. DesRoches, E. G. Campbell, S. R. Rao, K. Donelan, T. G. Ferris, A. Jha, R. Kaushal, D. E. Levy, S. Rosenbaum, A. E. Shields, D. Blumenthal, Electronic health records in ambulatory care—a national survey of physicians. *N. Engl. J. Med.* **359**, 50–60 (2008).
- U. S. Congress, American Recovery and Reinvestment Act, 2009.
- D. Blumenthal, Stimulating the adoption of health information technology. *N. Engl. J. Med.* **360**, 1477–1479 (2009).
- S. Shea, G. Hripcsak, Accelerating the use of electronic health records in physician practices. *N. Engl. J. Med.* **362**, 192–195 (2010).
- Meaningful use criteria—final rule; <http://edocket.access.gpo.gov/2010/pdf/2010-17207.pdf> [accessed 8 October 2010].
- D. Blumenthal, M. Tavenner, The “meaningful use” regulation for electronic health records. *N. Engl. J. Med.* **363**, 501–504 (2010).
- J. J. Nadler, G. J. Downing, Liberating health data for clinical research applications. *Sci. Transl. Med.* **2**, 18cm6 (2010).
- G. M. Church, Genomes for all. *Sci. Am.* **294**, 46–54 (2006).
- W. Burke, B. M. Psaty, Personalized medicine in the era of genomics. *JAMA* **298**, 1682–1684 (2007).
- D. A. Cortese, A vision of individualized medicine in the context of global health. *Clin. Pharmacol. Ther.* **82**, 491–493 (2007).
- G. S. Ginsburg, H. F. Willard, Genomic and personalized medicine: Foundations and applications. *Transl. Res.* **154**, 277–287 (2009).
- B. J. Edwards, C. Haynes, M. A. Levenstien, S. J. Finch, D. Gordon, Power and sample size calculations in the presence of phenotype errors for case/control genetic association studies. *BMC Genet.* **6**, 18 (2005).
- Electronic Medical Records and Genomics (eMERGE); <https://www.gwas.net/> [accessed 8 October 2010].
- S. Murphy, S. Churchill, L. Bry, H. Chueh, S. Weiss, R. Lazarus, Q. Zeng, A. Dubey, V. Gainer, M. Mendis, J. Glaser, I. Kohane, Instrumenting the health care enterprise for discovery research in the genomic era. *Genome Res.* **19**, 1675–1681 (2009).
- C. Friedman, P. O. Alderson, J. H. Austin, J. J. Cimino, S. B. Johnson, A general natural-language text processor for clinical radiology. *J. Am. Med. Inform. Assoc.* **1**, 161–174 (1994).
- J. C. Denny, A. Spickard III, K. B. Johnson, N. B. Peterson, J. F. Peterson, R. A. Miller, Evaluation of a method to identify and categorize section headers in clinical documents. *J. Am. Med. Inform. Assoc.* **16**, 806–815 (2009).
- G. K. Savova, P. V. Ogren, P. H. Duffy, J. D. Buntrock, C. G. Chute, Mayo Clinic NLP system for patient smoking status identification. *J. Am. Med. Inform. Assoc.* **15**, 25–28 (2008).
- H. Xu, S. P. Stenner, S. Doan, K. B. Johnson, L. R. Waitman, J. C. Denny, MedEx: A medication information extraction system for clinical narratives. *J. Am. Med. Inform. Assoc.* **17**, 19–24 (2010).
- R. A. Wilke, R. L. Berg, P. Peissig, T. Kitchner, B. Sijercic, C. A. McCarty, D. J. McCarty, Use of an electronic medical record for the identification of research subjects with diabetes mellitus. *Clin. Med. Res.* **5**, 1–7 (2007).
- J. C. Denny, R. A. Miller, L. R. Waitman, M. A. Arrieta, J. F. Peterson, Identifying QT prolongation from ECG impressions using a general-purpose natural language processor. *Int. J. Med. Inform.* **78** (Suppl. 1), S34–S42 (2009).
- C. A. McCarty, P. Peissig, M. D. Caldwell, R. A. Wilke, The Marshfield Clinic Personalized Medicine Research Project: 2008 scientific update and lessons learned in the first 6 years. *Personalized Med.* **5**, 529–542 (2008).
- D. M. Roden, J. M. Pulley, M. A. Basford, G. R. Bernard, E. W. Clayton, J. R. Balser, D. R. Mays, Development of a large-scale de-identified DNA biobank to enable personalized medicine. *Clin. Pharmacol. Ther.* **84**, 362–369 (2008).
- The NUGene Project, <http://www.nugene.org/> [accessed 8 October 2010].
- W. A. Kukull, R. Higdon, J. D. Bowen, W. C. McCormick, L. Teri, G. D. Schellenberg, G. van Belle, L. Jolley, E. B. Larson, Dementia and Alzheimer disease incidence: A prospective cohort study. *Arch. Neurol.* **59**, 1737–1746 (2002).
- I. J. Kullo, J. Fan, J. Pathak, G. K. Savova, Z. Ali, C. G. Chute, Leveraging informatics for genetic studies: Use of the electronic medical record to enable a genome-wide association study of peripheral arterial disease. *J. Am. Med. Inform. Assoc.* **17**, 568–574 (2010).
- M. D. Ritchie, J. C. Denny, D. C. Crawford, A. H. Ramirez, J. B. Weiner, J. M. Pulley, M. A. Basford, K. Brown-Gentry, J. R. Balser, D. R. Mays, J. L. Haines, D. M. Roden, Robust replication of genotype-phenotype associations across multiple diseases in an electronic medical record. *Am. J. Hum. Genet.* **86**, 560–572 (2010).
- J. C. Denny, M. D. Ritchie, D. C. Crawford, J. S. Schildcrout, A. H. Ramirez, J. M. Pulley, M. A. Basford, D. R. Mays, J. L. Haines, D. M. Roden, Identification of genomic predictors of atrioventricular conduction: Using electronic medical records as a tool for genome science. *Circulation* **122**, 2016–2021 (2010).
- Open Health Natural Language Processing (OHNLP) Consortium, <http://www.ohnlp.org/>
- J. C. Denny, J. F. Peterson, N. N. Choma, H. Xu, R. A. Miller, L. Bastarache, N. B. Peterson, Extracting timing and status descriptors for colonoscopy testing from electronic medical records. *J. Am. Med. Inform. Assoc.* **17**, 383–388 (2010).
- J. Friedlin, S. Grannis, J. M. Overhage, Using natural language processing to improve accuracy of automated notifiable disease reporting. *AMIA Annu. Symp. Proc.* 207–211 (2008).
- T. J. Love, T. Cai, E. W. Karlson, Validation of psoriatic arthritis diagnoses in electronic medical records using natural language processing. *Semin. Arthritis Rheum.* **40**, 413–420 (2011).
- K. P. Liao, T. Cai, V. Gainer, S. Goryachev, Q. Zeng-treitler, S. Raychaudhuri, P. Szolovits, S. Churchill, S. Murphy, I. Kohane, E. W. Karlson, R. M. Plenge, Electronic medical records for discovery research in rheumatoid arthritis. *Arthritis Care Res.* **62**, 1120–1127 (2010).
- Provisional guidance on the implementation of the 1997 standards for federal data on race and ethnicity; <http://minorityhealth.hhs.gov/templates/browse.aspx?lvl=2&lvlID=172> [accessed 8 October 2010].
- M. K. Wynia, S. L. Ivey, R. Hasnain-Wynia, Collection of data on patients' race and ethnic group by physician practices. *N. Engl. J. Med.* **362**, 846–850 (2010).

40. A. L. Price, N. J. Patterson, R. M. Plenge, M. E. Weinblatt, N. A. Shadick, D. Reich, Principal components analysis corrects for stratification in genome-wide association studies. *Nat. Genet.* **38**, 904–909 (2006).
41. G. B. Melton, N. Raman, E. S. Chen, I. N. Sarkar, S. Pakhomov, R. D. Madoff, Evaluation of family history information within clinical documents and adequacy of HL7 clinical statement and clinical genomics family history models for its representation: A case report. *J. Am. Med. Inform. Assoc.* **17**, 337–340 (2010).
42. W. G. Feero, M. B. Bigley, K. M. Brinner; Family Health History Multi-Stakeholder Workgroup of the American Health Information Community, New standards and enhanced utility for family health history information in the electronic health record: An update from the American Health Information Community's Family Health History Multi-Stakeholder Workgroup. *J. Am. Med. Inform. Assoc.* **15**, 723–728 (2008).
43. Surgeon General's Family Health History Initiative; <http://www.hhs.gov/familyhistory/> [accessed 8 October 2010].
44. R. W. Allen, M. H. Criqui, A. V. Diez Roux, M. Allison, S. Shea, R. Detrano, L. Sheppard, N. D. Wong, K. H. Stukovsky, J. D. Kaufman, Fine particulate matter air pollution, proximity to traffic, and aortic atherosclerosis. *Epidemiology* **20**, 254–264 (2009).
45. A. V. Diez-Roux, On genes, individuals, society, and epidemiology. *Am. J. Epidemiol.* **148**, 1027–1032 (1998).
46. A. V. Diez Roux, S. S. Merkin, D. Arnett, L. Chambless, M. Massing, F. J. Nieto, P. Sorlie, M. Szklo, H. A. Tyroler, R. L. Watson, Neighborhood of residence and incidence of coronary heart disease. *N. Engl. J. Med.* **345**, 99–106 (2001).
47. M. S. Mujahid, A. V. Diez Roux, J. D. Morenoff, T. E. Raghunathan, R. S. Cooper, H. Ni, S. Shea, Neighborhood characteristics and hypertension. *Epidemiology* **19**, 590–598 (2008).
48. C. J. Patel, J. Bhattacharya, A. J. Butte, An environment-wide association study (EWAS) on type 2 diabetes mellitus. *PLoS One* **5**, e10746 (2010).
49. J. S. Schildcrout, M. A. Basford, J. M. Pulley, D. R. Masys, D. M. Roden, D. Wang, C. G. Chute, I. J. Kullo, D. Carrell, P. Peissig, A. Kho, J. C. Denny, An analytical approach to characterize morbidity profile dissimilarity between distinct cohorts using electronic medical records. *J. Biomed. Inform.* **43**, 914–923 (2010).
50. Consolidated Health Informatics Initiative; <http://www.hhs.gov/healthit/chiinitiative.html> [accessed 4 October 2010].
51. Logical Observations Identifiers Names and Codes (LOINC), <http://loinc.org/> [accessed 8 October 2010].
52. RxNorm; <http://www.nlm.nih.gov/research/umls/rxnorm/> [accessed 12 October 2010].
53. J. C. Breitner, S. J. Haneuse, R. Walker, S. Dublin, P. K. Crane, S. L. Gray, E. B. Larson, Risk of dementia and AD with prior exposure to NSAIDs in an elderly community-based cohort. *Neurology* **72**, 1899–1905 (2009).
54. E. B. Larson, L. Wang, J. D. Bowen, W. C. McCormick, L. Teri, P. Crane, W. Kukull, Exercise is associated with reduced risk for incident dementia among persons 65 years of age and older. *Ann. Intern. Med.* **144**, 73–81 (2006).
55. **Funding:** The eMERGE Network was initiated and funded by the National Human Genome Research Institute, with additional funding from the National Institute of General Medical Sciences through grants U01-HG-004610 (Group Health Cooperative), U01-HG-004608 (Marshfield Clinic), U01-HG-04599 (Mayo Clinic), U01HG004609 (Northwestern University), and U01-HG-04603 (Vanderbilt University, also serving as the Coordinating Center), and the State of Washington Life Sciences Discovery Fund award to the Northwest Institute of Medical Genetics. The Vanderbilt BioVU and the Synthetic Derivative were supported in part by Clinical and Translational Research Award grant 1 UL1 RR024975 from the National Center for Research Resources, NIH. Funding for the Northwestern Enterprise Data Warehouse (EDW) was supported in part by Clinical and Translational Research grant UL1RR025741 from the National Center for Research Resources, NIH. **Author contributions:** All authors participated in the design and interpretation of the experiments and results. A.N.K., J.A.P., P.L.P., L.R., K.M.N., N.W., P.K.C., J.P., C.G.C., S.J.B., R.L.C., and J.C.D. participated in the acquisition and analysis of data. A.N.K., J.A.P., P.L.P., C.G.C., K.M.N., N.W., I.J.K., J.C.D., and P.K.C. performed statistical analysis. A.N.K., P.L.P., K.M.N., J.C.D., and C.G.C. led data collection and validation from each participating site. All authors contributed toward writing and editing the manuscript. **Competing interests:** The authors declare that they have no competing interests.

Submitted 14 October 2010

Accepted 1 April 2011

Published 20 April 2011

10.1126/scitranslmed.3001807

Citation: A. N. Kho, J. A. Pacheco, P. L. Peissig, L. Rasmussen, K. M. Newton, N. Weston, P. K. Crane, J. Pathak, C. G. Chute, S. J. Bielinski, I. J. Kullo, R. Li, T. A. Manolio, R. L. Chisholm, J. C. Denny, Electronic medical records for genetic research: Results of the eMERGE consortium. *Sci. Transl. Med.* **3**, 79re1 (2011).