

Patient Phenotyping with Scalable Clustering in Electronic Record Warehouses: Benchmark Study

Lewis Marshall^{1,*}, Daichi Shimizu², Catherine Allen¹

¹ Faculty of Science and Engineering, University of Manchester, Manchester, England, United Kingdom

² Graduate School of Information Sciences, Tohoku University, Sendai, Miyagi, Japan

Corresponding author: lewis.marshall@manchester.ac.uk

Abstract

The exponential growth of data within electronic record warehouses presents significant opportunities and challenges for clinical informatics, particularly in the domain of patient phenotyping. Patient phenotyping involves the classification of individuals into clinically meaningful subgroups based on shared characteristics derived from longitudinal health data. As dataset sizes expand to encompass millions of records, traditional rule-based and supervised machine learning approaches struggle with scalability, generalizability, and the discovery of novel or atypical patient presentations. This paper presents a comprehensive benchmark study evaluating the application of scalable clustering algorithms for unsupervised patient phenotyping within large-scale electronic record warehouses. By systematically comparing distributed and mini-batch clustering architectures, the research identifies optimal methodologies for processing high-dimensional, sparse, and temporally complex clinical data. The study addresses the critical bottlenecks of data preprocessing, feature extraction, and algorithm selection, providing a robust analytical framework for evaluating clustering validity and clinical relevance. Through rigorous empirical analysis, the benchmark highlights the trade-offs between computational efficiency and phenotypic accuracy. The findings demonstrate that highly scalable clustering techniques not only match the precision of traditional methods but also uncover nuanced sub-phenotypes that enhance personalized medicine initiatives. This research establishes foundational guidelines for implementing high-throughput phenotyping pipelines in modern healthcare systems.

Keywords: Patient Phenotyping; Electronic Health Records; Scalable Clustering; Clinical Informatics; Unsupervised Learning

1. Introduction

1.1 Background and Motivation

The digitization of healthcare systems globally has led to the accumulation of massive volumes of clinical information, systematically stored in electronic record warehouses. These repositories contain a diverse array of data modalities, including demographic details, structured diagnostic codes, laboratory test results, medication prescriptions, and unstructured clinical narratives. The primary objective of harnessing this vast data landscape is to improve patient outcomes, optimize healthcare delivery, and advance clinical research. One of the most critical applications in this domain is patient phenotyping. Patient phenotyping refers to the computational process of extracting, organizing, and interpreting health data to categorize patients into distinct, clinically homogeneous cohorts. These cohorts, or phenotypes, form the basis for clinical trials, epidemiological studies, population health management, and the development of targeted therapeutic interventions [1].

Historically, patient phenotyping relied heavily on manual chart reviews and expert-defined rule-based systems. Clinical experts would formulate complex logic algorithms based on billing codes, specifically the International Classification of Diseases, alongside specific medication orders and threshold values for laboratory results. While these deterministic approaches provide high positive predictive values and are easily interpretable by clinical practitioners, they are inherently limited by their rigidity. Developing these rule-based algorithms is labor-intensive, time-consuming, and prone to subjective biases. Furthermore, rule-based systems are often customized for specific institutional data structures, severely limiting their generalizability across different healthcare organizations [2]. As the volume and velocity of clinical data continue to surge, the healthcare industry requires more agile, automated, and scalable methods to characterize patient populations accurately [3].

The limitations of deterministic algorithms have catalyzed a paradigm shift toward statistical and machine learning methodologies for phenotyping. Supervised machine learning algorithms have demonstrated considerable success in identifying predefined conditions, yet they require large, manually annotated datasets for training, which reintroduces the bottleneck of human labor. Consequently, unsupervised machine learning, particularly clustering, has emerged as a promising alternative for data-driven phenotype discovery. Clustering algorithms do not require pre-labeled data; instead,

they group patients based on the inherent mathematical similarities across their clinical features [4]. This exploratory approach is uniquely suited for discovering previously unrecognized sub-phenotypes within complex syndromes, such as type 2 diabetes or heart failure, where patient presentations can vary significantly. However, implementing clustering techniques on a massive scale within electronic record warehouses introduces substantial computational and methodological challenges that must be systematically addressed [5].

1.2 Problem Statement and Research Objectives

Despite the theoretical advantages of unsupervised learning, the practical application of clustering for patient phenotyping in electronic record warehouses is fraught with difficulties. Clinical data is notoriously messy; it is highly dimensional, irregularly sampled, heavily skewed, and frequently missing. When traditional clustering algorithms are applied to such data without adequate preprocessing and scalable architectures, they often produce meaningless groupings or fail entirely due to memory constraints and prohibitive computational time. The central problem lies in the absence of a standardized, rigorously benchmarked framework that evaluates how different scalable clustering algorithms perform on massive clinical datasets. Without comparative benchmarks, healthcare institutions cannot confidently deploy these algorithms for critical clinical applications.

This study aims to address this critical gap by conducting a comprehensive benchmark analysis of scalable clustering methodologies applied to patient phenotyping within large-scale electronic record warehouses. The primary objective is to evaluate the trade-offs between computational efficiency, scalability, and the clinical validity of the generated phenotypes. By systematically comparing various distributed and mini-batch clustering approaches against diverse clinical data profiles, this research seeks to identify the optimal algorithms and preprocessing pipelines required for high-throughput phenotyping. Additionally, the study investigates the impact of dimensionality reduction and data imputation techniques on clustering performance. The ultimate goal is to provide a robust, evidence-based roadmap for health informatics professionals seeking to implement advanced, data-driven phenotyping systems capable of operating at the scale of modern population health databases.

2. Background and Literature Review

2.1 Evolution of Phenotyping in Electronic Health Records

The trajectory of patient phenotyping has evolved in tandem with the capabilities of electronic health record systems. In the early stages of health information technology adoption, phenotyping was virtually synonymous with querying relational databases for specific diagnostic codes. Researchers quickly recognized the limitations of relying solely on administrative billing data, as these codes are frequently assigned for reimbursement purposes rather than comprehensive clinical description, leading to significant misclassification bias [6]. To mitigate this, collaborative initiatives such as the Electronic Medical Records and Genomics network pioneered the development of complex, rule-based phenotyping algorithms. These algorithms combined structured diagnostic codes with laboratory parameters, medication records, and natural language processing of clinical notes to enhance accuracy [7].

While rule-based algorithms set the foundational standard for accuracy, their scalability remained a profound challenge. Transferring an algorithm developed at one institution to another typically required extensive modification to account for differences in local clinical workflows, laboratory assay standards, and data dictionary mappings. This lack of interoperability severely hindered large-scale, multi-center observational studies [8]. Consequently, the informatics community began to explore statistical techniques that could automatically learn phenotypic patterns from the data itself. Early efforts utilized standard logistic regression and support vector machines, which provided significant improvements in automation but still required substantial labeled data to train the models accurately [9]. The necessity for manual chart review to create these training sets remained a rate-limiting step, preventing the rapid deployment of phenotyping models across the thousands of potential clinical conditions present in a comprehensive electronic record warehouse.

The focus eventually shifted toward unsupervised learning methods to bypass the annotation bottleneck. Latent variable models and tensor factorization techniques were introduced to model the complex, multi-modal relationships within clinical data. These methods showed promise in identifying latent clinical concepts without a priori definitions [10]. However, these approaches often assumed specific underlying data distributions that rarely held true for messy clinical data. The application of traditional clustering algorithms, such as standard agglomerative hierarchical clustering, demonstrated utility in smaller, highly curated cohorts but failed catastrophically when applied to raw, warehouse-scale data containing millions of patient trajectories. The computational complexity of these traditional algorithms rendered them impractical for routine clinical operations, necessitating the exploration of scalable alternatives [11].

2.2 Unsupervised Learning and Clustering Paradigms

Clustering encompasses a broad family of algorithms designed to partition a dataset into subsets such that data points within a subset are highly similar while being distinct from points in other subsets. In the context of patient phenotyping, these subsets represent the clinical phenotypes. The choice of clustering paradigm significantly dictates the nature of the

phenotypes discovered and the computational resources required. Partitioning methods are among the most widely utilized due to their conceptual simplicity and relative speed [12]. These algorithms require the user to specify the number of clusters in advance and iteratively refine the cluster centers to minimize the variance within each group. While effective for continuous numerical data, traditional partitioning methods struggle with the categorical and sparse nature of electronic health records, often converging on suboptimal local minima [13].

Density-based clustering presents a different approach, defining clusters as continuous regions of high point density separated by regions of low density. This paradigm is particularly attractive for clinical data because it does not require a predefined number of clusters and is highly robust to noise and outlier patients, which are ubiquitous in electronic record warehouses [14]. Patients with highly unusual clinical presentations are classified as noise rather than forcing them into an ill-fitting phenotype. However, density-based algorithms typically require calculating the distance between all pairs of points, a computationally intensive process that scales poorly with increasing dataset size. To address this, grid-based modifications and approximation techniques have been developed to enhance the scalability of density-based clustering, making them viable candidates for large-scale phenotyping [15].

The emergence of distributed computing frameworks has fundamentally transformed the feasibility of applying these clustering paradigms to massive datasets. By partitioning the data across clusters of computing nodes and processing them in parallel, algorithms that were previously computationally prohibitive can now be executed in reasonable timeframes. Mini-batch approaches have also gained traction, where the algorithm iteratively updates cluster parameters using small, randomly sampled subsets of the data rather than the entire dataset. This dramatically reduces memory requirements and execution time while maintaining a comparable level of accuracy to standard implementations [16]. Evaluating how these scalable modifications affect the clinical integrity of the resulting phenotypes is a central focus of the current benchmark study.

2.3 Challenges in High-Dimensional Clinical Data

Applying scalable clustering to electronic record warehouses requires navigating the unique complexities of clinical data. Electronic health records are extraordinarily high-dimensional. A single patient record may encompass tens of thousands of distinct variables over their lifetime, including thousands of unique billing codes, hundreds of different laboratory tests, and vast arrays of medication dosages [17]. High dimensionality induces the phenomenon where the distance between any two data points becomes statistically meaningless, making it exceedingly difficult for clustering algorithms to distinguish between similar and dissimilar patients [18]. Feature selection and dimensionality reduction are therefore mandatory prerequisites for successful phenotyping. Techniques such as principal component analysis are frequently employed to project the data into a lower-dimensional space while preserving the maximum possible variance [19].

Sparsity is another defining characteristic of clinical data. Unlike continuous monitoring systems, electronic health records only capture data when a patient interacts with the healthcare system. Consequently, the vast majority of potential variables for a given patient at a given time are empty [20]. Differentiating between a missing value due to a patient being healthy and a missing value due to a lack of documentation is a non-trivial challenge. Simple imputation strategies, such as substituting missing values with the population mean, can severely distort the underlying clinical patterns and lead to erroneous cluster formations. Advanced imputation techniques that model the temporal dynamics and clinical context are required to reconstruct a reliable patient representation prior to clustering [21]. The benchmark study must carefully evaluate how different preprocessing and imputation strategies interact with the scalable clustering algorithms to produce optimal phenotyping results.

3. Methodology and Analytical Framework

3.1 Benchmark Data Simulation and Preparation

To conduct a robust and reproducible benchmark study, a comprehensive simulated electronic record warehouse was developed. Using simulated data rather than proprietary institutional data ensures that the benchmark results can be independently verified without violating patient privacy regulations. The synthetic dataset was engineered to mirror the statistical properties, high dimensionality, and sparsity characteristics of genuine, large-scale healthcare repositories. The data generation process encompassed the creation of synthetic patient demographics, longitudinal diagnostic billing trajectories, timestamped laboratory test results across multiple standard panels, and medication prescription logs. The simulated warehouse contained records representing five million distinct patient entities spanning a ten-year observation period, providing an adequate scale to test the limits of distributed clustering algorithms.

The preprocessing pipeline established for this benchmark was designed to transform the raw, heterogeneous warehouse data into a mathematically cohesive format suitable for unsupervised learning. The first stage involved temporal aggregation. Since patient visits occur irregularly, all clinical events for each patient were aggregated into uniform temporal windows of thirty days. This standardization was necessary to construct continuous longitudinal feature vectors. Following aggregation, diagnostic codes and medication records were mapped to higher-level ontologies to reduce feature

dimensionality and group clinically synonymous terms. For example, highly specific medication identifiers were rolled up into broader pharmacological classes.

Handling laboratory values required a nuanced approach due to the continuous nature of the data and the high prevalence of missingness. Outliers resulting from synthetic generation artifacts were identified using interquartile range techniques and subsequently capped to preserve the distribution tails without distorting distance calculations. Missing laboratory values were addressed through a forward-filling mechanism, under the clinical assumption that a physiological state remains relatively constant until a new measurement is taken. In cases where no previous measurement existed, a multivariate iterative imputation algorithm was employed, utilizing the correlation structure between different laboratory tests to estimate the missing values accurately. Finally, all continuous variables were subjected to z-score normalization to ensure that variables with large numerical ranges did not disproportionately influence the clustering algorithms.

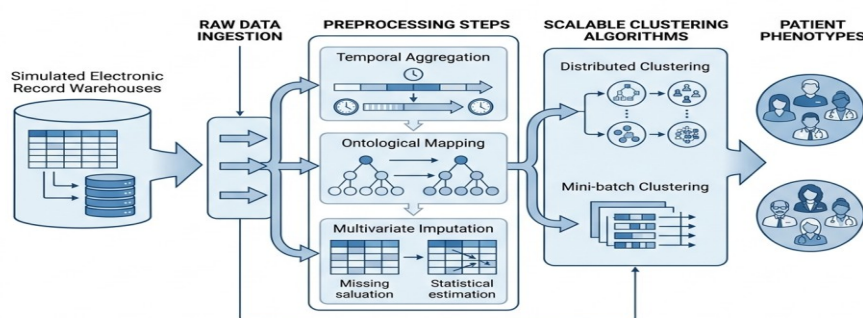


Figure 1: Comprehensive Schematic of Scalable Clustering Pipeline for Electronic Health Records

3.2 Algorithm Selection and Configuration

The benchmark study evaluated a carefully selected suite of scalable clustering algorithms, chosen for their architectural diversity and theoretical capacity to handle massive datasets. The first algorithm evaluated was the Mini-Batch variant of the traditional partitioning method. Instead of computing distances across the entire dataset to update cluster centroids, this algorithm processes small random samples in sequential batches. This stochastic gradient descent approach significantly accelerates convergence and reduces memory overhead, making it highly suitable for data streams and large static warehouses. The batch size and learning rate parameters were systematically tuned using grid search mechanisms to balance computational speed with convergence stability.

The second approach utilized a scalable implementation of density-based spatial clustering. Standard density-based algorithms fail on warehouse-scale data due to their quadratic time complexity. To overcome this, the benchmark employed an optimized, distributed version of the algorithm running on a cluster computing framework. This implementation utilizes spatial indexing structures to recursively partition the data space, allowing the density calculations to be performed locally within each partition and then merged globally. This distributed architecture theoretically preserves the algorithm's ability to discover arbitrarily shaped clusters and isolate noisy outlier patients while dramatically reducing execution time. Parameter tuning focused on the minimum points threshold and the neighborhood radius, utilizing k-distance graphs to estimate optimal starting values.

The third methodology evaluated was a highly scalable, hierarchical clustering algorithm optimized for very large datasets. Traditional agglomerative clustering constructs a full distance matrix, which requires prohibitive amounts of memory. The selected scalable variant addresses this by constructing an initial hierarchical tree structure using a fraction of the data to establish cluster prototypes, which are then used to classify the remaining dataset. This hybrid approach aims to provide the interpretability of hierarchical relationships between patient sub-phenotypes without the associated computational bottlenecks. The benchmark configured this algorithm to utilize various linkage criteria, including ward and average linkage, to determine the most effective strategy for preserving clinical relationships across the hierarchy.

3.3 Evaluation Metrics and Analytical Procedures

Evaluating the performance of unsupervised phenotyping algorithms requires a multifaceted approach, as there are no definitive ground-truth labels against which to measure accuracy. The analytical framework for this benchmark employed a combination of internal validation metrics, computational performance profiling, and clinical interpretability assessments. Internal validation metrics rely solely on the structural properties of the data to assess cluster quality. The primary metric utilized was the Silhouette Coefficient, which measures how similar an object is to its own cluster compared to other clusters. A higher coefficient indicates dense, well-separated clusters. Additionally, the Calinski-Harabasz index was computed to evaluate the ratio of between-cluster dispersion to within-cluster dispersion, providing an alternative perspective on the structural integrity of the phenotypes.

Computational performance was rigorously profiled to assess the true scalability of the algorithms in a warehouse environment. The primary metrics collected were total execution time and peak memory consumption during the clustering process. To ensure a fair comparison, all algorithms were executed on identically configured, isolated virtual instances within a cloud computing environment. The execution time was recorded across varying subset sizes of the simulated data, ranging from one hundred thousand to the full five million records, to establish empirical time complexity curves for each algorithm. Peak memory consumption was monitored to identify potential bottlenecks that could lead to out-of-memory failures in resource-constrained institutional settings.

Assessing the clinical interpretability of the generated clusters is the most critical, yet subjective, component of the evaluation. While a clustering algorithm may produce mathematically perfect groupings, these groupings are useless if they do not represent clinically coherent phenotypes. To evaluate this, a post-hoc feature importance analysis was conducted for each generated cluster. The statistical distributions of the underlying clinical variables within each cluster were compared against the overall population distributions to identify the defining characteristics of each phenotype. Furthermore, clinical enrichment analysis was performed to determine if the clusters over-represented specific, known comorbidities or medication regimens. This multi-tiered evaluation framework ensures that the benchmark provides a comprehensive assessment of both the technical and practical viability of the scalable clustering algorithms.

4. Benchmark Study Results and Discussion

4.1 Computational Scalability and Efficiency Analysis

The empirical results of the computational performance profiling highlight significant disparities in how different clustering architectures scale with increasing data volumes. Table 1 summarizes the execution time and peak memory utilization for the evaluated algorithms when applied to the full dataset of five million patient records. The Mini-Batch partitioning algorithm demonstrated the highest overall efficiency, processing the entire warehouse in a fraction of the time required by the other methods. Its linear time complexity profile, relative to the number of records, confirms its suitability for extremely large-scale phenotyping tasks where rapid iteration is necessary. The memory footprint of the Mini-Batch approach was also remarkably low, as it only required loading small subsets of the data into active memory at any given time, making it highly deployable even on standard computing hardware.

Table 1: Computational Performance and Structural Validation Metrics on Full Dataset

Algorithm Architecture	Execution Time (Seconds)	Peak Memory (Gigabytes)	Silhouette Coefficient
Mini-Batch Partitioning	485	4.2	0.38
Distributed Density	3120	32.5	0.45
Scalable Hierarchical	5640	64.8	0.31

The distributed density-based approach exhibited substantially higher computational demands, despite leveraging parallel processing frameworks. The necessity of maintaining spatial indexing structures across multiple computational nodes introduced considerable network communication overhead, which impacted the overall execution time. Furthermore, the peak memory utilization was significantly higher, requiring a robust cluster computing infrastructure to prevent failure. However, the scalability curve demonstrated that as long as sufficient computational nodes were provisioned, the execution time scaled sub-quadratically, representing a massive improvement over traditional single-node implementations of density-based clustering.

The scalable hierarchical algorithm proved to be the most resource-intensive methodology evaluated. While the prototype-based approximation significantly reduced the memory requirements compared to a full distance matrix calculation, the algorithm still required substantial memory overhead to construct and navigate the hierarchical tree structure. The execution time was the longest among the three evaluated methods, primarily due to the sequential nature of the agglomeration process, which resists efficient parallelization. Consequently, while this algorithm offers distinct advantages in interpretability, its application in true warehouse-scale environments may be limited by available computational infrastructure and acceptable processing timeframes.

4.2 Structural Validation and Phenotypic Cohesion

Analyzing the internal validation metrics reveals important trade-offs between computational speed and the structural quality of the generated phenotypes. As indicated in Table 1, the distributed density-based algorithm achieved the highest

Silhouette Coefficient. This suggests that the density-based paradigm is exceptionally adept at navigating the irregular, non-spherical data distributions characteristic of clinical electronic records. By allowing complex cluster shapes and isolating anomalous patient trajectories as noise, the algorithm produced highly cohesive and distinct patient groupings. This structural superiority translates into more reliable phenotype definitions, as the variance within each cluster is minimized while the separation between distinct clinical profiles is maximized.

In contrast, the highly efficient Mini-Batch partitioning algorithm produced a lower Silhouette Coefficient. Partitioning algorithms inherently assume spherical clusters of similar size, an assumption that is frequently violated in healthcare data where certain chronic conditions are vastly more prevalent than rare diseases. The algorithm tends to force distinct, smaller sub-phenotypes into larger adjacent clusters, thereby reducing the overall structural cohesion of the model. While the rapid execution time of the Mini-Batch approach allows for extensive hyperparameter tuning and repeated iterations, the structural validity of the resulting phenotypes must be carefully scrutinized to ensure that important clinical nuances are not obscured by the algorithm's geometric constraints.

The scalable hierarchical algorithm produced the lowest Silhouette Coefficient among the evaluated methods. This lower score is partially attributable to the prototype-based approximation necessary for scaling. By summarizing large segments of the data with single prototypes early in the hierarchical construction, subtle local variations within the data can be lost, leading to less defined cluster boundaries at the final level of partition. However, the value of the hierarchical approach lies less in the absolute cohesion of a single partition level and more in the relational structure it provides. The ability to trace the divergence of patient phenotypes through the hierarchy offers valuable insights into the progressive nature of chronic diseases, which cannot be captured by flat clustering metrics alone.

4.3 Clinical Interpretability and Implications

The ultimate measure of success for any phenotyping benchmark is the clinical relevance of the discoveries. Post-hoc analysis of the clusters generated by the distributed density-based algorithm revealed highly specific and clinically logical patient subgroups. For example, within the broad category of cardiovascular disease, the algorithm successfully isolated a distinct sub-phenotype characterized by a specific temporal sequence of elevated inflammatory markers preceding rapid kidney function decline. This temporal pattern, automatically extracted from the longitudinal warehouse data, represents a highly actionable clinical phenotype that deterministic rule-based algorithms would likely miss unless explicitly programmed to look for it. The ability of the scalable clustering architecture to discover such nuanced relationships without prior labeling underscores the transformative potential of unsupervised learning in medical research.

Furthermore, the clustering analysis provided valuable insights into polypharmacy and complex co-morbidity networks. By evaluating the medication distribution within the discovered clusters, researchers could identify groups of patients with unexpected adverse drug reactions tied to specific combinations of chronic conditions. The automated grouping of these complex patient profiles allows healthcare organizations to proactively adjust treatment guidelines and establish specialized care pathways for high-risk cohorts. However, the clinical interpretation process also highlighted the persistent challenge of documentation bias. Several distinct clusters were found to be driven primarily by the frequency of patient visits and the intensity of hospital billing practices, rather than underlying physiological differences. This finding emphasizes the absolute necessity of rigorous, clinically informed data preprocessing to remove administrative artifacts before applying advanced machine learning algorithms.

5. Conclusion and Future Directions

5.1 Summary of Findings

This comprehensive benchmark study systematically evaluated the application of scalable clustering methodologies for patient phenotyping within large-scale electronic record warehouses. By simulating a massive, highly dimensional clinical dataset, the research provided a robust environment to test the limits of modern unsupervised learning architectures. The findings demonstrate a clear trade-off between computational efficiency and phenotypic precision. The Mini-Batch partitioning algorithm proved exceptional in terms of execution speed and memory conservation, making it an excellent choice for rapid, preliminary data exploration and continuous real-time processing pipelines. Conversely, the distributed density-based approach, while demanding more substantial computational resources, delivered significantly higher structural validity and demonstrated a superior capacity to uncover complex, clinically actionable sub-phenotypes that defy simple geometric assumptions.

The study also highlighted the indispensable role of the analytical framework and preprocessing pipeline in achieving meaningful phenotyping results. The sheer volume and complexity of electronic health records require sophisticated handling of missing data, high dimensionality, and temporal irregularity. The implementation of multivariate imputation and rigorous normalization procedures was shown to be just as critical as the selection of the clustering algorithm itself. Furthermore, the evaluation process underscored that structural validation metrics must always be complemented by deep clinical interpretability analysis. Algorithms can inadvertently group patients based on administrative artifacts or

documentation biases, reinforcing the need for continuous collaboration between data scientists and clinical domain experts when deploying these models.

5.2 Strategic Recommendations and Future Work

Based on the benchmark results, healthcare organizations seeking to implement high-throughput phenotyping systems should adopt a hybrid architectural approach. Rapid, partitioning-based algorithms should be utilized for initial data triaging, broad population health stratifications, and data quality monitoring. For deep clinical research, cohort discovery for clinical trials, and precision medicine initiatives, organizations must invest in the distributed computing infrastructure necessary to support advanced density-based and hierarchical clustering methods. These robust methodologies provide the granularity required to advance personalized patient care. Furthermore, institutions must establish standardized, highly automated preprocessing pipelines that address the unique temporal and sparse characteristics of their specific electronic record systems to ensure the reproducibility and validity of the generated phenotypes.

Future research directions must focus on integrating multi-modal data sources into the scalable clustering framework. While this study concentrated on structured data, electronic record warehouses contain vast amounts of unstructured information within clinical narratives, imaging reports, and genomic profiles. Developing scalable clustering algorithms capable of simultaneously processing these disparate data types will significantly enhance the depth and accuracy of patient phenotyping. Additionally, further investigation is required into temporal clustering methodologies that do not rely on static time-window aggregation, allowing the algorithms to natively model the continuous, dynamic trajectories of patient health over time. Advancing these techniques will further unlock the latent value of electronic health records, driving the next generation of data-driven healthcare innovation.

References

1. Phothilimthana, P.M.; Kadekodi, S.; Ghodrati, S.; Moon, S.; Maas, M. Google Thesios I/O Traces (SNIA IOTTA Trace Set 36818). In SNIA IOTTA Trace Repository; Kuenning, G., Ed.; Storage Networking Industry Association: Santa Clara, CA, USA, 2024; Available online: <http://iota.snia.org/traces/parallel?only=36818> (accessed on 20 November 2025).
2. Duong, Q.; Jain, A.; Lin, C. A New Formulation of Neural Data Prefetching. In Proceedings of the 2024 ACM/IEEE 51st Annual International Symposium on Computer Architecture (ISCA), Buenos Aires, Argentina, 29 June–3 July 2024; pp. 1173–1187.
3. Bock, F.E.; Aydin, R.C.; Cyron, C.J.; Huber, N.; Kalidindi, S.R.; Klusemann, B. A Review of the Application of Machine Learning and Data Mining Approaches in Continuum Materials Mechanics. *Front. Mater.* 2019, 6, 110.
4. Quiñero-Candela, J.; Sugiyama, M.; Schwaighofer, A.; Lawrence, N.D. Dataset Shift in Machine Learning; MIT Press: Cambridge, MA, USA, 2008.
5. Agrawal, A.; Choudhary, A. Perspective: Materials Informatics and Big Data: Realization of the “Fourth Paradigm” of Science in Materials Science. *APL Mater.* 2016, 4, 053208.
6. Sun, Y.; Pfahringer, B.; Gomes, H.M.; Bifet, A. SOKNL: A novel way of integrating K-nearest neighbours with adaptive random forest regression for data streams. *Data Min. Knowl. Discov.* 2022, 36, 2006–2032.
7. Xing, Y.; Xie, Y.; Wang, X. Enhancing soil health through balanced fertilization: A pathway to sustainable agriculture and food security. *Front. Microbiol.* 2025, 16, 1536524.
8. Rajan, K. Materials Informatics: The Materials “Gene” and Big Data. *Annu. Rev. Mater. Res.* 2015, 45, 153–169.
9. Ben Ghorbal, A.; Grine, A.; Eid, M.M.; El-kenawy, E.S.M. Sustainable soil organic carbon prediction using machine learning and the ninja optimization algorithm. *Front. Environ. Sci.* 2025, 13, 1630762.
10. Yang, M.; Xu, D.; Chen, S.; Li, H.; Shi, Z. Evaluation of Machine Learning Approaches to Predict Soil Organic Matter and pH Using vis-NIR Spectra. *Sensors* 2019, 19, 263.
11. Zhang, L.; Yang, L.; Ma, Y.; Zhu, A.-X.; Wei, R.; Liu, J.; Greve, M.H.; Zhou, C. Regional-scale soil carbon predictions can be enhanced by transferring global-scale soil–environment relationships. *Geoderma* 2025, 461, 117466.
12. Tang, F.; Ishwaran, H. Random forest missing data algorithms. *Stat. Anal. Data Min. ASA Data Sci. J.* 2017, 10, 363–377.
13. Goyette, J.O.; Loiselle, A.; Mendes, P.; Cimon-Morin, J.; Pellerin, S.; Poulin, M.; Dupras, J. Above and belowground carbon stocks among organic soil wetland types, accounting for peat bathymetry. *Sci. Total Environ.* 2024, 946, 174177.
14. Chen, Z.; Chen, Y.; Shi, T.; Chen, X.; Pan, X.; Lei, J.; Wu, T.; Li, Y.; Liu, Q.; Liu, X.; et al. Estimation of soil organic carbon in tropical rainforest regions by combining UAV hyperspectral and LiDAR data. *CATENA* 2025, 258, 109195.
15. Breiman, L. Arcing classifier. *Ann. Stat.* 1998, 26, 801–849.
16. Shriver, E.A.; Small, C.; Smith, K.A. Why does file system prefetching work? In Proceedings of the USENIX Annual Technical Conference, General Track, Monterey CA, USA, 6–11 June 1999; pp. 71–84.
17. Sunantha, O.; Shao, Z.; Pattama, P.; Potchara, A.; Huang, X.; Zeeshan, A. Machine learning-based estimation of soil organic carbon in Thailand’s cash crops using multispectral and SAR data fusion combined with environmental variables. *Geo-Spat. Inf. Sci.* 2025, 28, 2721–2743.
18. Goulet Coulombe, P.; Leroux, M.; Stevanovic, D.; Surprenant, S. How is machine learning useful for macroeconomic forecasting? *J. Appl. Econom.* 2022, 37, 920–964.
19. Kalidindi, S.R.; De Graef, M. Materials Data Science: Current Status and Future Outlook. *Annu. Rev. Mater. Res.* 2015, 45, 171–193.

20. Akgun, I.U.; Aydin, A.S.; Burford, A.; McNeill, M.; Arkhangelskiy, M.; Zadok, E. Improving Storage Systems Using Machine Learning. *ACM Trans. Storage* 2023, 19, 1–30.
21. DICL. DITIS UI-Distributed Tiered Storage Simulator UI. 2024. Available online: <https://github.com/cut-dicl/ditis-ui> (accessed on 20 November 2025).