



Smart Feature Engineering for Large Health Datasets

Benjamin Clark
Asst Professor

Abstract

A rapidly proliferating corpus of clinical research harnessing the power of machine learning has substantial implications for healthcare feature engineering. As a broad umbrella encompassing data preprocessing, quality control, transformation, and generation, feature engineering addresses a major bottleneck in the production of predictive models. Automated feature engineering systems are increasingly deployed at scale to meet the challenges of generating the vast quantity of predictive features necessary for successful, generalizable, and clinically useful machine learning systems. Such clinical feature engineering systems produce features that are applied in a predictive setting after the fact and not explicitly linked to clinical care, but nevertheless involve substantial risk. A principled examination of a clinical feature engineering system can be framed in terms of six core components: data governance; compliance with privacy and regulatory constraints; appropriate validation and prospective evaluation; consideration of data biases; the use of effective data-quality and preprocessing pipelines; and sound candidate feature generation, scoring, and selection strategies. Health systems typically possess an assemblage of rich and diverse, yet underutilized, information with the potential to contribute meaningfully to clinical prediction problems. Electronic health record (EHR) data, comprising clinical notes, laboratory values, medication orders, and procedure codes; over a decade's worth of length and width dataset and point-of-care laboratory test results; continuous biosensor measurements; DNA sequencing data; biomarkers derived from imaging; and drug compounds targeting genotypes provide raw material for hundreds of prediction problems in diverse specialties. However, machine learning in healthcare exhibits a stunning lack of reproducibility: many predictive models fail to retain their accuracy in different cohorts, and those that do are seldom incorporated into routine clinical care. A significant bottleneck underlying this failure lies with the feature engineering step.

Keywords: Clinical Feature Engineering, Automated Feature Engineering Systems, Healthcare Machine Learning, Predictive Model Pipelines, Electronic Health Records (EHR), Multimodal Clinical Data, Data Governance in Healthcare AI, Privacy and Regulatory Compliance, Prospective Model Validation, Bias-Aware Feature Design, Data Quality and Preprocessing Pipelines, Candidate Feature Generation and Selection, Feature Scoring Strategies, Reproducible Clinical ML, Biosensor and Time-Series Data, Imaging-Derived Biomarkers, Genomic and Sequencing Features, Point-of-Care Laboratory Data, Clinical Model Generalization, Production-Grade Healthcare AI.



1. Introduction

Feature engineering is a challenging, yet critical, part of the supervised machine learning pipeline. While there are many exciting, new ML models that have demonstrated their success in other fields, such as language, vision, or games, they rely on large, curated datasets that require a careful selection, crafting, and engineering of features. Following a similar approach in healthcare is very difficult, due mainly to the many different sources of data that are used, the fact that they originate from different hospitals and that they must be updated each time a new prospective study is proposed. Because of these challenges, such studies can take months or even years before a model has been built, tested, and validated. Ideally, careful consideration of the features and data sources should counterbalance that, improving predictive performance, enhancing generalization, or leading to the discovery of new risk factors. At least from an algorithmic perspective, it would be sufficient to automate feature engineering and treat it as a black box. Automating the generation of predictive features in healthcare then pushes these methods one step further.

Healthcare feature engineering is the process of gathering, selecting, and preparing features from one or more healthcare data sources, in particular Electronic Health Records (EHR) databases. Such data sources contain a wealth of information on patients, including their medical history, medications, and diagnostics, which can be used to predict future clinical outcomes. However, these sources vary from site to site and evolve over time, introducing heterogeneity and uncertainty into the dataset. Therefore, feature generation consists mainly of three stages: generating a source of feature candidates by applying domain knowledge and heuristics, deploying data preparation and cleaning pipelines to ensure that the features produced can then be used safely in a future model, and scoring the candidates with an automated or manual process to finally select the best-ranked features. Past success of the first stage in using domain knowledge and heuristics to mine predictive subgraphs with temporal constraints has demonstrated the potential of feature generation methods in healthcare.

1.1. Overview of Healthcare Feature Engineering

Healthcare feature engineering encompasses the design of features from large, heterogeneous, longitudinal clinical data sources—particularly electronic health records, but also imaging,

genomic, sensor, and textual data—to support clinical decision-making.

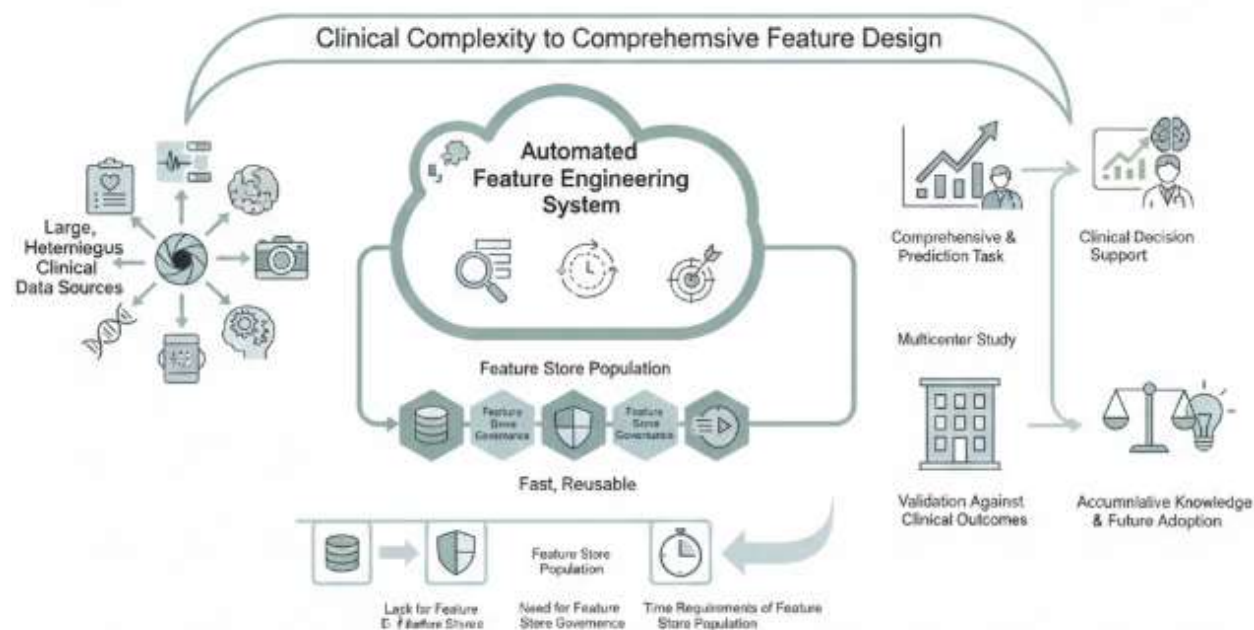


Fig 1: Scaling Clinical Intelligence: Automated Feature Engineering and Governed Feature Stores for Multimodal, Reproducible Healthcare AI

It aims to produce a comprehensive set of clinically relevant, time-varying features for a targeted patient population that maximally inform the prediction task, integrating all data sources available for those patients. Clinical complexity makes comprehensive feature design highly challenging, motivating automated feature engineering systems that mirror the pipeline of data scientists with rich domain knowledge.

Two features or sets of features are rarely used across studies, limiting the accumulative knowledge-building typical in research. While proposed feature sets are often extensive, they typically focus on a single domain (e.g., EHR diseases, meds, lab results) without standardized integration across domains. The reasons are threefold: lack of a feature store, the need for feature store governance, and the time requirements of feature store population. An automated feature engineering system fulfills these needs and enables fast, comprehensive, and reusable feature engineering for a target study population. Applying such a system in a multicenter study



provides insights on usage patterns and enables validation against clinical outcomes, supporting future adoption of broader, prospective feature design pipelines.

2. Foundations of Feature Engineering in Healthcare

Considerable amounts of semi-structured, unstructured, and high-dimensional medical data have been amassed over the past few decades. Prominent data sources include electronic health records (EHRs), medical imaging and diagnostic reports, genomic sequencing data, biosensors, mobile health applications, and wearable sensors. These modalities can complement clinical data, provide additional insights about a patient's condition, or even replace clinical assessments such as the ejection fraction score for heart disease. Despite their apparent utility in clinical decision-making, these auxiliary datasets are still seldomly used in supervised machine learning tasks.

Within the context of supervised machine learning, a feature is defined as a measurable characteristic of a phenomenon being observed. Feature stores serve as centralized repositories for storing and serving data features for use in machine learning models. These features are engineered and made available for use in supervised learning tasks, with properties that allow for easy discovery, management, and use across projects. Features store metadata also provide a source of potential value by tracking data provenance. Well-maintained feature stores can significantly reduce the time and effort to develop ML models and increase reproducibility by allowing for easy retrieval of the same feature set.

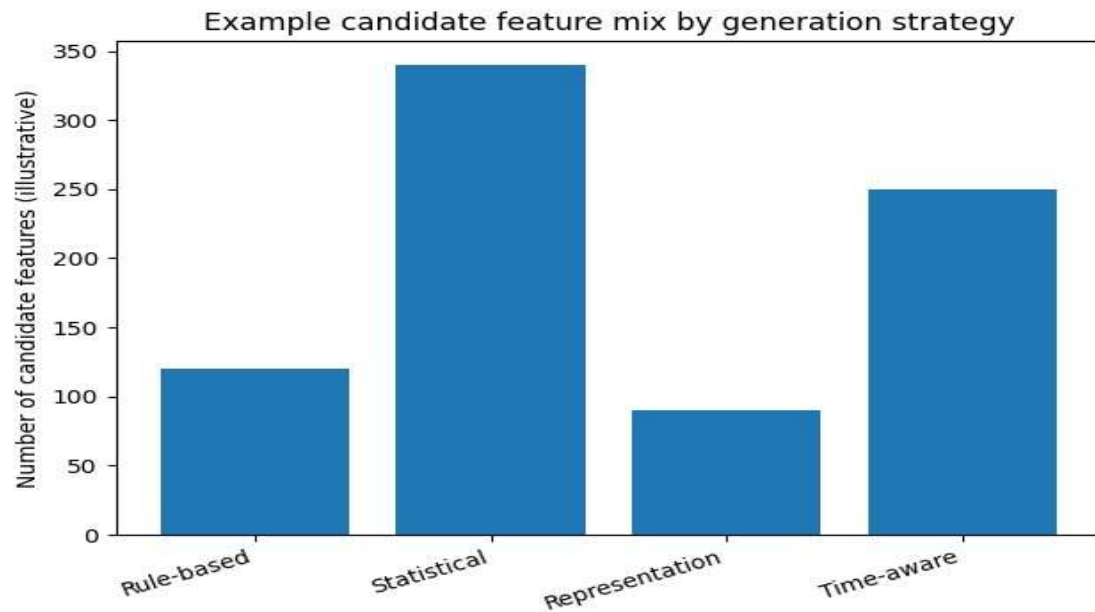
2.1. Clinical relevance and data modalities

Healthcare machine learning relies on diverse clinical datasets and requires features with medical interpretations that are relevant to clinical decision-making. The most common data modalities in healthcare are electronic health records, medical images, genomic data, and data from wearable and stationary sensors. The data sources are inherently different in their content and the dimensions they measure which is further complicated by the vast amount of heterogeneous information that is aggregated across time. In addition, most of the data sources available for ML are generated during the normal course of clinical practice and when other scientific investigations take place which raises challenges for data quality. Proper mappings must enable the application of ML techniques on the collected data sources.

In healthcare environments, data sources are most frequently accessed in keep-it-simple strategies, whether through the data lake/exploitation or the hospital's multiple data products.



Authors focus on the first strategy, where a clinic uses the data lake and the other clinics make use of the data distributed in a subsequent event. When taking this approach, data quality is still an issue and authors combined information cleaning and data harmonization cycle, applied on retrospective healthcare ML applications, with validation in cross-site frameworks. The feature engineering process must then take the complete process into account but be divided in two distinct machines: one for data quality assurance and the second for feature generation.



Equation 1) Data quality & preprocessing equations (clean → normalize → impute → deduplicate → harmonize)

1.1 Cleaning as constraint projection (validity rules)

Let raw measurement be x . Suppose “plausible clinical bounds” for a lab are $[L, U]$. A standard cleaning operation is clipping (winsorization):

1. If $x < L$, set $x' = L$
2. If $L \leq x \leq U$, set $x' = x$



3. If $x > U$, set $x' = U$

Compactly:

$$x' = \min(\max(x, L), U)$$

1.2 Normalization

(A) Min-max normalization (map to $[0,1]$)

Given a feature column $x_1, \dots, x_n$:

1. Compute $x_{\min} = \min_i x_i$, $x_{\max} = \max_i x_i$
2. For each value:

$$z_i = \frac{x_i - x_{\min}}{x_{\max} - x_{\min}}$$

(B) Z-score standardization (mean 0, variance 1)

1. Mean:

$$\mu = \frac{1}{n} \sum_{i=1}^n x_i$$

2. Population std:

$$\sigma = \sqrt{\frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2}$$

3. Standardize:

$$z_i = \frac{x_i - \mu}{\sigma}$$



1.3 Imputation (missing values)

Let x_i be missing for some indices.

Median imputation

1. Compute median $m = \text{median}(\{x_i: x_i \text{ observed}\})$
2. Replace missing values:

$$x_i^{\text{imp}} = \begin{cases} x_i, & x_i \text{ observed} \\ m, & x_i \text{ missing} \end{cases}$$

Mean imputation is the same form with $m \leftarrow \mu$.

Linear interpolation (for time series at times t_k)

If $x(t_a)$ and $x(t_b)$ are observed and $t_a < t < t_b$ is missing:

$$x(t) = x(t_a) + \frac{t - t_a}{t_b - t_a} (x(t_b) - x(t_a))$$

2.2. Defining features, features stores, and lineage

A feature is a single measurable property or characteristic of a phenomenon. In a healthcare context, clinical features refer to attributes or constructs captured through diagnostic or prognostic procedures, exam findings, or tests. Laboratory test results represent features measured at a point in time. Genomic expression analysis examines hundreds of thousands of features simultaneously, each referring to expression levels of specific genes. Medical imaging is fundamentally about measuring features, and different imaging modalities extract different features relevant for specific disease processes. A feature store is a repository for features, usually organized by their predictive task. Feature stores facilitate discovery and reuse. Automated Feature Engineering in a large healthcare environment implies resilient and scalable software products. Features are gathered, modeled, stored, and tested for algorithmic fitness in a systemic way.

Maintaining feature lineage—tracking how features were created and the precise data (origin, version, state) and processing pipelines used to generate them—greatly aids collaboration, improves reproducibility, addresses quality issues, and enhances compliance with regulatory requirements. Lineage tracking typically follows a hierarchical approach: features



combine datasets, and datasets depend on raw data and other derived datasets. The entire history of the computation underlying a clinical feature can be captured and stored with metadata about the feature. Individual nodes in the lineage can be attributed and versioned, can record the health of the lineage, and can be used to trigger de-duplication, re-creation, and recalibration. Individual features may have governance models.

patient_id	time_hr	creatinine_after_dedup	age
101	0	1.1	55
102	0		70
103	0	0.6	40
104	0	9.8	62
105	0	1.4	33

3. Automated Feature Engineering: Techniques and Architectures

Automated feature engineering encompasses four broad categories: rule-based generation, statistical feature generation, representation learning, and the generation of time-aware features. Each category may apply to a single feature or groups of features. Rule-based features may be generated with any suitable method (e.g., user or crowd-sourced rules, templates, philosophy-based ontologies). Statistical feature generation constitutes the application of functions to selected groups of features across any user-defined window (a.k.a., moving or sliding window). Representation learning encapsulates the suite of (semi-)unsupervised techniques for learning representations from large data sets. Finally, time-aware features are those whose values depend explicitly on time, areas of time, or other temporal characteristics.

With a library of candidate regions of interest generated by detection, a data provenance strategy designed to facilitate tracking and scoring candidate region features, and a functional pipeline in place to clean and prepare the region of interest features, the remaining open questions may be framed along two dimensions: the data quality aspects of preprocessing and the implementation of a pipeline capable of effecting (1) common operations such as input data cleaning, normalization, de-duplication, and temporal alignment across sites, and (2) specific target definitions calling for advanced harmonization (i.e., creation of features incorporating data from more than one site).

3.1. Feature generation strategies

A variety of techniques may be applied to automate feature generation from healthcare data. Literature identifies methods that can be categorized as rule-based approaches, statistical methods of feature generation, data-centric representation learning methods that facilitate the construction of task-specific features directly from the data, and time-aware feature generation techniques. Regardless of the specific techniques utilized, candidates must be scored—distinguishing between features worth pursuit and those that should be automatically pruned—and scored candidates subsequently selected.

Conceptually, healthcare features are derived from the processing of raw data: either a transformation is applied to the underlying data or a set of other features is used to train a predictive model for a new feature. Data preprocessing methods, including cleaning, normalization, imputation, de-duplication, harmonization across different data-collection sites, and domain-specific preprocessing workflows (e.g., image and natural-language processing), can be implemented in a standalone data-quality pipeline prior to candidate feature engineering.

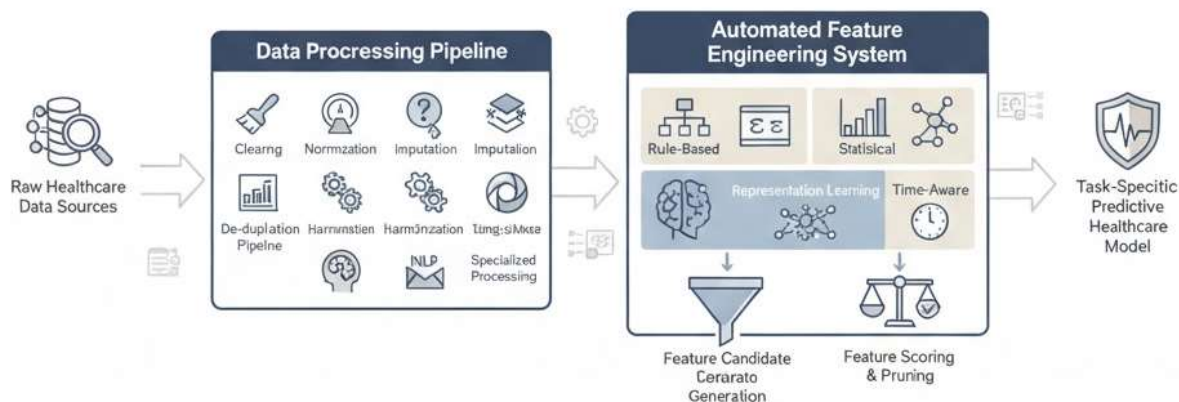


Fig 2: Automated Feature Engineering Systems in Healthcare: A Multimodal Pipeline for Candidate Generation, Scoring, and Clinical Validation

3.2. Data quality and preprocessing pipelines

A second pipeline addresses data quality and feature preprocessing tasks. Cleaning removes invalid, improbable, and contradictory values. Normalization scales values to a common range while preserving semantic meaning. Imputation fills in missing values through median/mode, value interpolation, predictive models, or generative methods. De-duplication detects and reconciles duplicate records for the same entity at the same point in time.



Harmonization reconciles semantically equivalent but contextually different records across data sources. Feature-specific preprocessing handles transformations, discretization, binning, embeddings, dimensionality reduction, and algorithm-specific transformations such as one-hot encoding for categorical variables.

Maintaining a Feature Store to centralize and standardize features while supporting lineage tracking increases reproducibility and transparency. A feature is an individual measurable characteristic or property of a phenomenon being observed. Automated feature engineering seeks to reduce the manual effort that data scientists spend on feature creation for machine learning models by developing tools, algorithms, and system support. Three automated steps are candidate generation, scoring, and selection.

4. Systems for Large-Scale Healthcare Environments

Data governance encompasses access control and logging procedures, de-identification methods (pseudonymization, k-anonymization, data masking), and assessment of trusted data-sharing frameworks (secure multi-party computation, trusted third-party-computation), allowing multiple parties to analyze sensitive data without revealing source information. Because clinical data often span multiple sites for many patients, harmonization needs to go beyond cleaning and normalizing the data to feature generation and candidate scoring processes. Data privacy and security must be firmly established to assure compliance with healthcare regulations. With reference to the United States, Australia, and Europe, regulatory oversight includes not only the Health Insurance Portability and Accountability Act and the General Data Protection Regulation but explicit links to Food and Drug Administration regulation of medical devices and drug trials, and considerations for limited patient consent and health-information-sharing agreements.

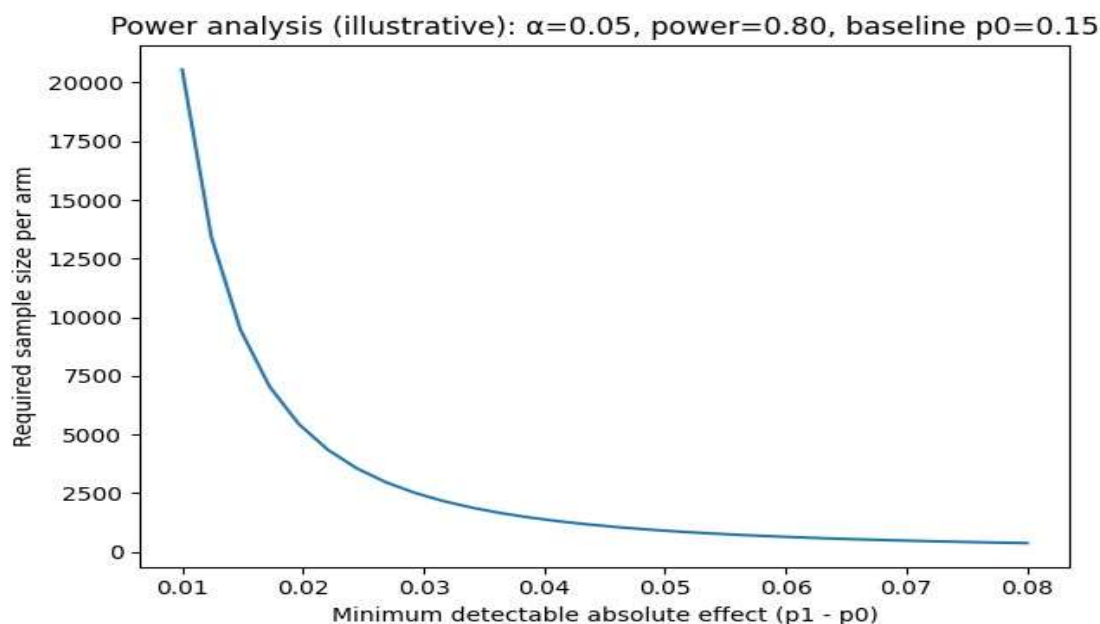
Evaluation and validation in retrospective studies focus on generalization across sites and populations, while prospective assessment establishes the safety and clinical effectiveness of a deployed system. Beyond receiver-operating-characteristic-area-under-the-curve scores, true clinical utility hinges on statistical power and, hence, the extent and timing of study recruitment. Post-validation management tasks—monitoring for adverse behavior, testing statistical hypotheses of changed performance, and executing planned cross-validation—respond to the open-endedness of clinical studies through exploratory analysis and pondering the root causes of any emerge performance shifts.

4.1. Data governance, privacy, and security



A comprehensive data governance framework assures that data remains secure, private, and confidential even when used across many institutions in a distributed, collaborative environment. Mechanisms for governing access control to data must be granular and fine-tuned to give the right user the appropriate levels of access for their current task. Audit logging keeps track of data usage; it allows tracking of the frequency and patterns of access, thereby revealing suspicious and potentially malicious activity. When raw clinical datasets contain sensitive personally identifiable information (PII), data must undergo de-identification before sharing. Sharing across sites using secure multi-party computation (MPC) allows complex joint analysis without disclosing PII, but it tends to be slow and users need specialized programming skills.

Healthcare data is sensitive by nature. Leaving unmonitored access to sensitive data across multiple external institutions is risky for the data owners, creating privacy and confidentiality issues. Compliance with healthcare privacy regulations, such as the Health Insurance Portability and Accountability Act (HIPAA) in the United States or the General Data Protection Regulation (GDPR) in Europe, is essential in a healthcare data governance plan. Data governance metrics involve identifying what data has been processed and how. Data retention policies describe how long each dataset should be kept and how it will be securely disposed of afterwards.





Equation 2) Time-aware & statistical feature generation (moving windows)

Let patient p have measurements $\{(t_j, x_j)\}$. For a window length W , define the set in the window ending at time t :

$$\mathcal{I}(t) = \{j: t - W < t_j \leq t\}$$

2.1 Window mean

$$\text{mean}_W(t) = \frac{1}{|\mathcal{I}(t)|} \sum_{j \in \mathcal{I}(t)} x_j$$

2.2 Window variance

1. Window mean $\bar{x}_W(t)$ as above
2. Then

$$\text{var}_W(t) = \frac{1}{|\mathcal{I}(t)|} \sum_{j \in \mathcal{I}(t)} (x_j - \bar{x}_W(t))^2$$

2.3 Trend (slope) in a window (linear regression)

Fit $x \approx \beta_0 + \beta_1 t$ for points in $\mathcal{I}(t)$.

Let window points be (t_j, x_j) for $j \in \mathcal{I}(t)$. Define:

$$\bar{t} = \frac{1}{m} \sum t_j, \quad \bar{x} = \frac{1}{m} \sum x_j, \quad m = |\mathcal{I}(t)|$$

Then the least squares slope is:

$$\beta_1 = \frac{\sum (t_j - \bar{t})(x_j - \bar{x})}{\sum (t_j - \bar{t})^2}$$

This produces a time-aware “trend” feature (e.g., creatinine rising over 48h).



4.2. Compliance with healthcare regulations

Frameworks mapping automated feature engineering into three classes of health care regulations include the HIPAA Privacy and Security Rules, the European Union General Data Protection Regulation (GDPR), and regulatory considerations from the U.S. Food and Drug Administration (FDA). Retention time is another important regulatory factor determining the frequency data are updated and the prospective use of a feature, such as when it is included in a data retention policy for automated and continuous feature operation. Consent, data use agreements, and contractual requirements from institutional review boards (IRBs) and endpoint-contributing data owners are considered for Trans-Union and Trans-Atlantic multimodality healthcare feature repositories.

feature	predictive_score	stability_across_sites	leakage_risk
mean_creatinine_24h	0.42	0.9	0.1
trend_creatinine_48h	0.38	0.76	0.15
med_count_7d	0.31	0.88	0.08
note_embedding_dim17	0.35	0.7	0.2
icu_bed_occupancy_site	0.12	0.55	0.05

5. Evaluation, Validation, and Deployment

Adequate evaluation is critical to determine if automated features are safe for clinical use and capable of performing as well as or better than hand-engineered alternatives. Retrospective validation of candidate features in electronic healthcare record data through offline modelling provides an initial assessment of practical utility. This validation can encompass standard validation techniques used in machine learning, including train-test splits, cross-site validation for assessing generalizability, performance calibration, and error analysis. Candidates are then deployed into prospective clinical settings in an automated manner with pipelines that monitor for scoring errors and track conciliation. A/B testing of predictive and prognostic models (and potentially risk stratification scores) in live clinical settings — as facilitated by the availability of a feature store — provides an opportunity to assess value of features, ideally against clinical endpoints such as in-hospital mortality, ICU admissions, or length of stay.



Ensure cross-site study population sizes are suitably large and that analyses can be stratified when necessary to cover potential differences in feature efficacy. Allocating sufficient time until results become available is critical. Ideally, deploy interventions with sufficiently low expected effect size, otherwise statistical power rapidly becomes overly demanding. When modelling the likely size of the intervention effect, consider the time-horizon and relevant data epochs. Upon deployment, design the statistical tests with the necessary α and β values, determine the minimum difference of interest, and calculate the necessary sample size. Ensure suitable monitoring mechanisms are in place to detect any anomalies after intervention. Finally, ascertain patient safety through trigger systems capable of identifying deleterious outcomes.

5.1. Offline validation and cross-site generalization

Worker performance, in terms of the AI model serving as a sub-component, must therefore be assessed in some detail. Macro-level validation targets for each pair of sites—e.g. a HF test set specifically for HF patients in sites A and B—must be established and monitored at deployment time. Critical metrics are calibrated, and candidate agents flagged according to their error analysis.

A standard external validation pipeline has also been established, including the generation of separate test sets from prospective cohorts (i.e. operating on data unseen by the learning agents). Initial preparations of the target cohorts, including de-identification, are performed by Site A. Evaluation agents are then built for the cohorts, with performance estimates returned to Site A. A predefined set of evaluation metrics is applied that enable dimension-based checks and considerations of medical validity for the evaluation results. A second pipeline applies the same checks for HT models, with agents routed through a model serving framework to facilitate the process.

At the other end, the computational dependencies of supervised learning models can also be scrutinized. By tracking feature use, test and train split colors are applied to pre-train/test cohorts, enabling a clear view of sites-only, sites-across, unseen, and test-test test sets. Cross-site validation by test split color further allows rapid scans of known strength for model theft and inclusion detection—enough to inform the use of success prediction systems.

Dominant logical mechanisms and hyper-parameter signatures within trained models can finally be parsed, enabling the recovery of hyponymy trees for an interactive model inspection tool. The visualization ensures that summarization can extend to modelling efforts across HT and HF elements.

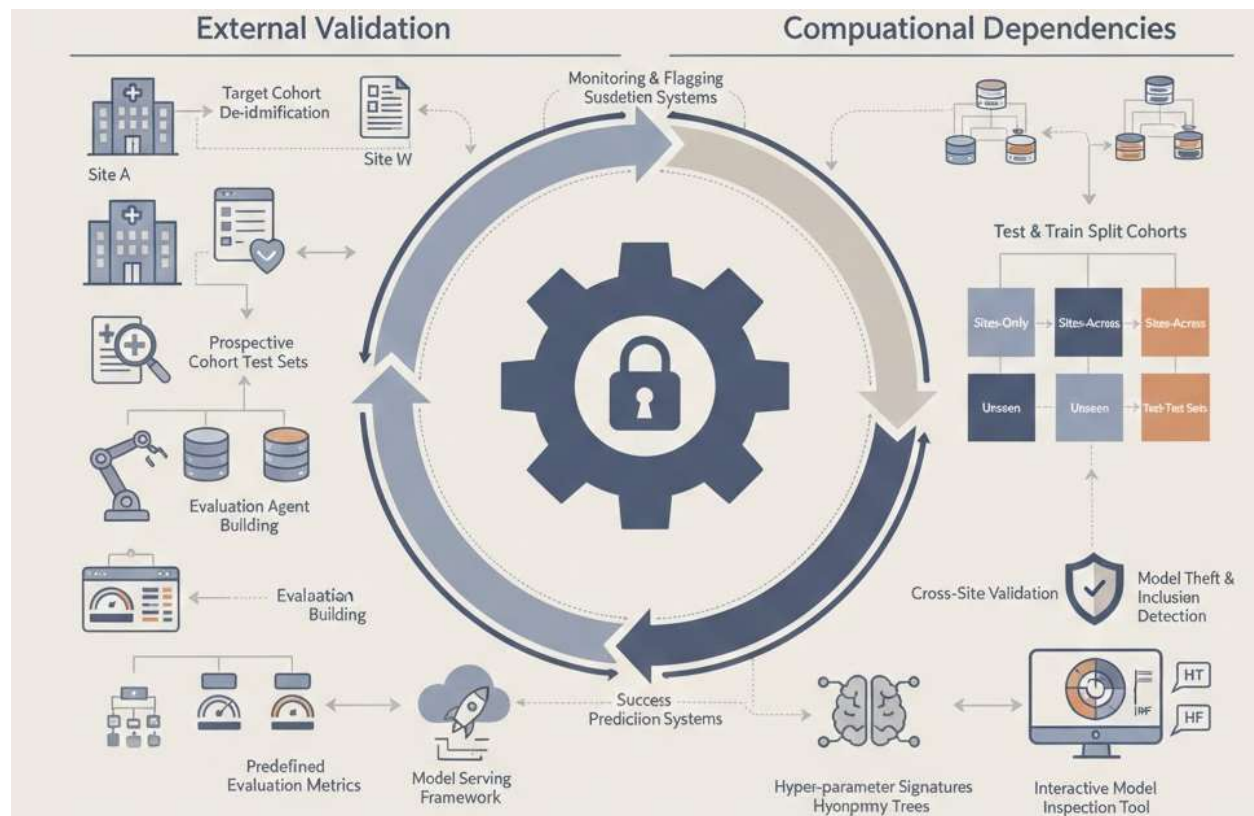


Fig 3: Cross-Site Clinical Validation: A Framework for Multi-Cohort Evaluation, Computational Dependency Tracking, and Interactive Model Inspection

5.2. Prospective evaluation and A/B testing in clinical settings

After offline validation, a deployment pipeline is used to move model-dependent features into a separate features store, followed by prospective evaluation in a clinical setting. Live traffic is assigned to candidate and baseline features following randomized assignment, with discernible clinical endpoints and an a priori power analysis determining sample size. Monitoring systems track commencement of the trial, and safety checks are established to stop the trial if the trial-determining feature fails to have a considerable effect in the direction signaled by the trials-determining model.



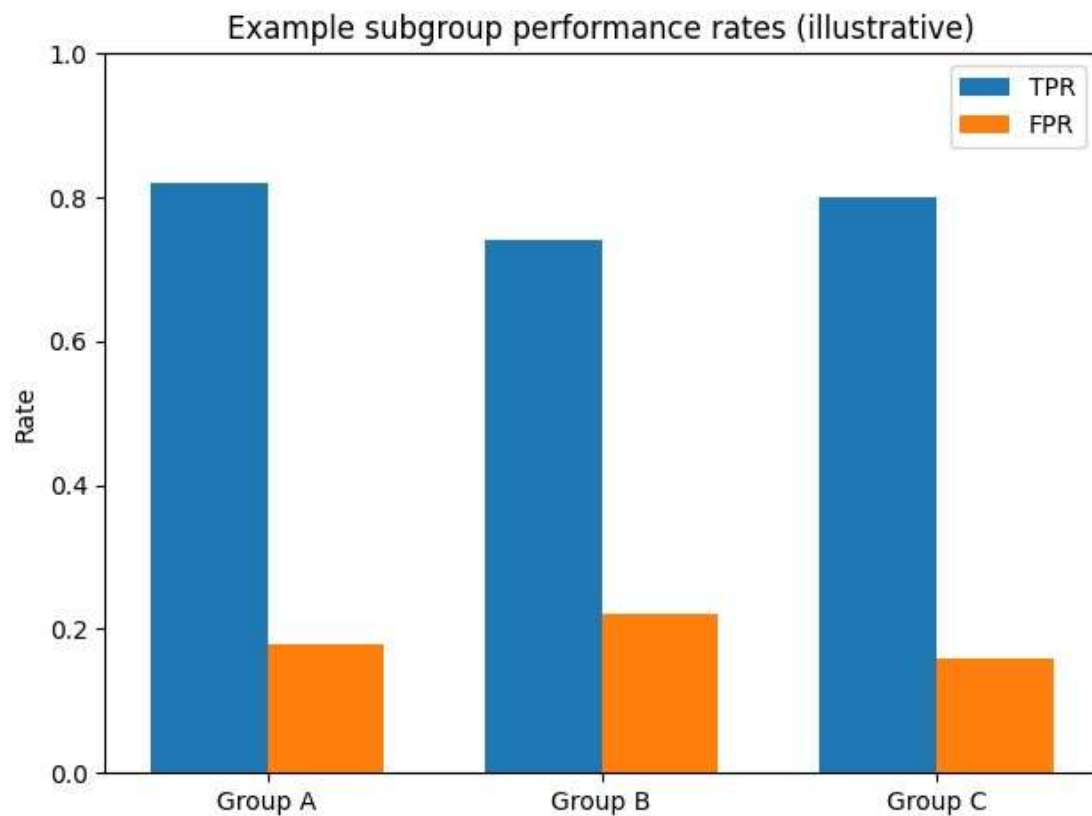
Prospective evaluation using A/B testing methodology is a natural next step for the validation of candidate features once offline requirements are met. This procedure permits the quantification of the clinical effect of the newly supplied model-dependent feature, retaining the benefits of assignment to features in production as feature-dependent active learning. The flexibility afforded by reduced data-collection costs, high accuracy, and substantial effort in debiasing of the model-determining features permits formal A/B testing.

6. Challenges and Risks in Automated Feature Engineering

Automated Feature Engineering Systems in Large-Scale Healthcare Data Environments: High-dimensional outcome-sensitive feature engineering critically influences the outcome of machine learning, yet few methods exist. In clinical settings, high-dimensional EHR and imaging data may contain features that are costly or risky to evaluate prospectively, necessitating automated feature engineering. Data governance, privacy, and security must satisfy industry standards; automated feature generation and risk mitigation tools should be tailored to healthcare's unique characteristics.

Data Bias, Fairness, and Equity Implications: Automated feature engineering can inadvertently produce problematic models that encode racial, gender, or economic bias. Systemically, analysis can estimate fairness-related model properties across protected attributes, examine performance on subgroups, and compute group-wise fairness metrics; remediation strategies, if needed, include sample stratification and in-processing adjustment. From a risk perspective, approaches in healthcare data have shown that building a clinical pathway for an outcome of interest allows testing for bias as well as links to social determinants of health, making bias, fairness, and equity implications visible.

Privacy-Preserving Feature Generation: Scalable healthcare feature engineering leverages data of unknown ownership via core components for privacy-sensitive training and feature generation. Federated learning, a distributed machine learning approach, trains models closest to the raw data, sidestepping raw data transfer and ownership concerns. Privacy accounting and differential privacy ensure model release and evaluation satisfy privacy requirements. Within an automated pipeline, potential feature generators are ranked by privacy risk, and resource-aware differential privacy correctly weights real and synthetic features. Reported privacy risk estimates promote stakeholder trust; clearly articulated costs lower adversarial investment.



Equation 3) Candidate feature scoring & selection (a basic, complete math formulation)

Let there be candidate features $f_1, \dots, f_M$. Define three common scoring ingredients:

3.1 Predictive strength (example: correlation with outcome)

For numeric outcome y , Pearson correlation:

1. Means: $\bar{f} = \frac{1}{n} \sum f_i, \bar{y} = \frac{1}{n} \sum y_i$
2. Covariance: $\text{cov}(f, y) = \frac{1}{n} \sum (f_i - \bar{f})(y_i - \bar{y})$
3. Std devs: σ_f, σ_y



4. Correlation:

$$r = \frac{\text{cov}(f, y)}{\sigma_f \sigma_y}$$

3.2 Cross-site stability (example: variance of effect across sites)

If you estimate a site-specific effect θ_s (e.g., AUC lift or coefficient), stability can be inverse dispersion:

$$\text{stability} = \frac{1}{1 + \text{Var}_s(\theta_s)}$$

3.3 Leakage / risk penalty

If leakage score $L(f) \in [0,1]$ (higher worse), then penalize.

3.4 Combined ranking score

A simple complete selection score:

$$S(f) = w_1 \cdot \text{predictive}(f) + w_2 \cdot \text{stability}(f) - w_3 \cdot \text{leakage}(f)$$

Select top- K :

$$\{f\}_{\text{selected}} = \arg \text{topK} S(f)_{f \in \{f_1, \dots, f_M\}}$$

6.1. Data bias, fairness, and equity implications

Multiple sources of bias may lead to datasets that fail to adequately represent entire segments of the healthcare population, creating issues of predictive fairness and inequity. From an equity perspective, it is critical to identify the potential sources and forms of bias in automated feature engineering systems, as well as their implications for population-specific predictions. Systematic methods, such as those by Barocas et al., can be employed to investigate and characterize bias in both the data and resulting deployed model. The key steps are as follows:



1. Consider important attribution questions: Who or what is the object of concern? (e.g., individuals, groups, models, predictions) Who or what is responsible for the outcome? (e.g., individuals, groups represented in the data, models) Who or what is affected by the outcome? (e.g., individuals, groups, society) Is the outcome rendered unjust? If so, how? (e.g., inequitable, unfair, unworthy)

2. Identify potential sources of bias in data: Bias can stem from issues relating to data generation, collection, integration, modeling/engineering, evaluation/selection, and actual deployment.

3. Diagnose and measure bias: Precision diagnostics reveal underrepresented subpopulations and their associated metrics, and fairness metrics expose whether the model merits concern from a fairness perspective.

4. Mitigate detected imbalance: Attention needs to be devoted to addressing detected concerns, possibly employing stratification by risk or tailoring the model in other ways.

6.2. Privacy-preserving feature generation

Privacy is paramount in healthcare, where sensitive patient information is collected, analyzed, and shared. The creation of features that rely on personal data must therefore be governed by data privacy and statistical disclosure risk principles to minimize the risk of patient identification. For cross-institutional sharing, privacy constraints often preclude direct data sharing and therefore necessitate special computational techniques that allow model training without direct access to underlying data. Federated learning, for example, enables decentralized ML wherein the algorithm is trained across multiple parties without exchanging the data used to build locally updated models. The final model is obtained by aggregating the locally trained models while meeting regulatory data use agreements. Other techniques, such as building a differentially private version of the features that aggregate information across a subset of users or designating a trusted data steward to manage the model training process, help ensure that direct or indirect identification is highly improbable. Therefore, a careful assessment of the risk of privacy violation and the definition of rigorous privacy budgets, together with differential privacy techniques, can facilitate privacy-compliant feature generation.

The importance of data privacy considerations increases when developing models capable of automated feature generation, such as GNNs or VAEs. Privacy-preserving training of these models has garnered increasing interest, driven by the success of ML and the associated



imperative of having access to large amounts of data. Previous work has proposed a federated learning methodology for training GNN models in a healthcare setting, with focus on monitoring heterogeneous disease progressions via subgraph-based learning. Additionally, several methods for privacy-preserving generative modeling using VAE architectures have recently been proposed. Thus, proper design considerations can enable automated feature generation for highly sensitive datasets.

7. Conclusion

Automated feature engineering for predictive models trained on healthcare data is crucial because of the associated economic and human costs of healthcare data science projects. Feature stores specifically designed for large clinical datasets enable the sharing, quality control, and systematic generation of data quality-enhancing features, rules, and cleaning procedures among data scientists and stakeholders. A mandated feature store addresses the challenge of ensuring that all services utilize optimal rules and features, but at the expense of introducing a bureaucratic overhead.

Herein, the central role of features in predictive models for clinical data pipelines, the techniques and architectures for automation, and the specialized requirements for large healthcare data environments are substantiated. Potential sources of bias during automated feature generation and the implications of such biases are examined. Frameworks for prospective validations in clinical settings that go beyond model performance and for the ethical generation of features given privacy concerns are explored. Finally, avenues for resolving open methodological questions and translating automated feature engineering systems into production are identified.



regulations; mechanisms for offline validation; and a framework for prospective evaluation and A/B testing in clinical settings.

Although Automated Feature Engineering holds considerable promise across multiple application domains, including healthcare, further exploration is necessary. For example, little is known about data bias, fairness, and equity implications; a novel set of heuristic fairness metrics has recently been proposed, yet practical approaches for bias mitigation remain scarce. Applying novel privacy-preserving techniques such as secure multiparty computation, federated learning, and differential privacy is another promising avenue, especially given heightened attention to data privacy and protection. Accelerating the application of Automated Feature Engineering in healthcare would ultimately enhance the safety, efficacy, and equity of predictive models, thereby addressing the pressing need for scalable solutions and enabling artificial intelligence to realize its full potential across healthcare and other domains.

References

1. Chen, I. Y., Joshi, S., Ghassemi, M., & Ranganath, R. (2021). Probabilistic machine learning for healthcare. *Annual Review of Biomedical Data Science*, 4, 393–415.
2. Kreimeyer, K., Dang, O., Spiker, J., Muñoz, M. A., Rosner, G., Ball, R., & Botsis, T. (2021). Feature engineering and machine learning for causality assessment in pharmacovigilance: Lessons learned from application to the FDA Adverse Event Reporting System. *Computers in Biology and Medicine*, 135, 104517.
3. Kolla, S. H. (2023). Large Language Model-Driven Enterprise Service Intelligence for Digital Workflow Transformation. *International Journal of Research and Applied Innovations*, 6(1), 8380-8391.
4. Garnica, O., Gómez, D., Ramos, V., Hidalgo, J. I., & Ruiz-Giardín, J. M. (2021). Diagnosing hospital bacteraemia in the framework of predictive, preventive and personalised medicine using electronic health records and machine learning classifiers. *EPMA Journal*, 12, 365–381.
5. Yoshida, Y., et al. (2021). Supervised machine learning-based prediction for in-hospital pressure injury development using electronic health records: A retrospective observational cohort study in a university hospital in Japan. *International Journal of Nursing Studies*, 119, 103932.
6. Wang, H., Wang, L., Lee, E. H., Zheng, J., Zhang, W., Halabi, S., Liu, C., Deng, K., Song, J., & Yeom, K. W. (2021). Decoding COVID-19 pneumonia: Comparison of deep learning and radiomics CT image signatures. *European Journal of Nuclear Medicine and Molecular Imaging*, 48, 1873–1882.



7. Bandi, V. D. V. K. (2023, November). Production-Grade Machine Learning Pipelines For Healthcare Predictive Analytics. *South Eastern European Journal of Public Health*, 189–205.
8. Yuan, J., Ran, X., Liu, K., Yao, C., Yao, Y., Wu, H., & Liu, Q. (2022). Machine learning applications on neuroimaging for diagnosis and prognosis of epilepsy: A review. *Journal of Neuroscience Methods*, 368, 109441.
9. Janjua, Z. H., Kerins, D., O'Flynn, B., & Tedesco, S. (2022). Knowledge-driven feature engineering to detect multiple symptoms using ambulatory blood pressure monitoring data. *Computer Methods and Programs in Biomedicine*, 217, 106638.
10. Kumaraswamy, N., Markey, M. K., Barner, J. C., & Rascati, K. (2022). Feature engineering to detect fraud using healthcare claims data. *Expert Systems with Applications*, 201, 118433.
11. Rajendran, R., & Karthi, A. (2022). Heart disease prediction using entropy based feature engineering and ensembling of machine learning classifiers. *Expert Systems with Applications*, 207, 117882.
12. Mangalampalli, B. M. Generative AI Applications In Healthcare Data Mart Design And Optimization.
13. Garriga, R., Mas, J., Abraha, S., Nolan, J., Harrison, O., Tadros, G., & Matic, A. (2022). Machine learning model to predict mental health crises from electronic health records. *Nature Medicine*, 28, 1240–1248.
14. Xie, F., Zhou, J., Lee, J. W., Tan, M., Li, S., Rajnithern, L. S., Chee, M. L., Chakraborty, B., Wong, A. I., Dagan, A., Ong, M. E. H., Gao, F., & Liu, N. (2022). Benchmarking emergency department prediction models with machine learning and public electronic health records. *Scientific Data*, 9, 658.
15. Sadasivuni, S., Saha, M., Bhatia, N., Banerjee, I., & Sanyal, A. (2022). Fusion of fully integrated analog machine learning classifier with electronic medical records for real-time prediction of sepsis onset. *Scientific Reports*, 12, 5711.
16. Davuluri, P. S. L. N. (2023). Integrating artificial intelligence into event-driven financial crime compliance platforms. *International Journal of Finance*, 36(6), 707-736.
17. Wu, Y., et al. (2022). Improved prediction of body mass index in real-world administrative healthcare claims databases. *Pharmacoepidemiology and Drug Safety*, 31, 1019–1028.
18. Yang, S., Varghese, P., Stephenson, E., Tu, K., & Gronsbell, J. (2023). Machine learning approaches for electronic health records phenotyping: A methodical review. *Journal of the American Medical Informatics Association*, 30(2), 367–381.
19. Sun, H., & Wang, X. (2023). High-dimensional feature selection in competing risks modeling: A stable approach using a split-and-merge ensemble algorithm. *Biometrical Journal*, 65(2), e2100164.
20. Zeng, D., et al. (2023). A general framework of nonparametric feature selection in high-dimensional data. *Biometrics*, 79(2), 951–963.



21. Yandamuri, U. S. (2023). An Intelligent Analytics Framework Combining Big Data and Machine Learning for Business Forecasting. Zenodo.
22. La Cava, W. G., et al. (2023). A flexible symbolic regression method for constructing interpretable clinical prediction models. *npj Digital Medicine*, 6, Article 116.
23. Ben Miled, Z., Dexter, P. R., Grout, R. W., & Boustani, M. (2023). Feature engineering from medical notes: A case study of dementia detection. *Heliyon*, 9(3), e14636.
24. Ben-Assuli, O., Heart, T., Klempfner, R., & Padman, R. (2023). Human-machine collaboration for feature selection and integration to improve congestive heart failure risk prediction. *Decision Support Systems*, 172, 113982.
25. Margeloiu, A., Simidjievski, N., Liò, P., & Jamnik, M. (2023). Weight predictor network with feature selection for small sample tabular biomedical data. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37(8), 9195–9203.
26. Bhattacharjee, A., Basak, S., & Kumari, P. (2023). A two-step feature selection procedure for relevant markers of squamous cell lung carcinoma using different survival models. *Healthcare Analytics*, 3, 100168.
27. Kolla, S. K., & Reddy, V. A. R. (2023). Deep Learning Architectures For Multimodal Medical Data Integration. *South Eastern European Journal of Public Health*, 248–260.
28. Wang, X., et al. (2023). Deep feature screening: Feature selection for ultra high-dimensional data via deep neural networks. *Neurocomputing*, 538, 126186.
29. Mukherjee, P., Humbert-Droz, M., Chen, J. H., & Gevaert, O. (2023). SCOPE: Predicting future diagnoses in office visits using electronic health records. *Scientific Reports*, 13, 11005.
30. Inness, C., et al. (2023). Predicting disease onset from electronic health records for population health management: A scalable and explainable deep learning approach. *Frontiers in Artificial Intelligence*, 6, 1287541.
31. Kumar, V. D. V. B. (2023). Automated feature engineering systems in large-scale healthcare data environments. *Journal of Neonatal Surgery*.