

Original Article

Intelligent Data Engineering and Generative AI for Biomarker Discovery in Neurodegenerative and Kidney Diseases

***Mallesham Goli**
Independent Researcher, USA.

Abstract:

Generative AI-Enhanced Data Engineering Pipelines for Predictive Biomarker Discovery in Alzheimer's Disease and Kidney Disease: an objective, evidence-based, formal study of methodologies, architectures, and implications, with clear, parsimonious argumentation and rigorous evaluation. Concise, objective synthesis of the study's aims, hypotheses, scope, and contributions; specification of research questions; expected impact on biomarker discovery. The clinical significance of Alzheimer's disease risk-modifying biomarkers is widely accepted. Nevertheless, despite pervasive data science activity in the search for predictive biomarkers, DNA-based predictors remain elusive, proteomic-based predictors are too often unreplicated, and AI-based predictors are often unvalidated and poorly understood. The soaring number of data repositories holds great potential for the discovery of predictive disease biomarkers; however, issues with data quality, integration, reproducibility, and lack of adequate engineering pipelines hinder this promise. Existing full data engineering pipelines are rarely employed. Generative AI is a novel, emerging area of research and application with potential to transform traditional information-technology and data-engineering infrastructure, with broad implications for data engineering for Alzheimer's disease, kidney disease, and the search for other predictive disease biomarkers. Generative AI is increasingly being used in the biomedical domain. Nevertheless, generative-AI-enhanced data-engineering pipelines that support the entire data-flow lifecycle of predictive biomarker discovery remain to be published. Questions include how generative AI can enhance pipelines, what data engineering contributions will be important to pipeline success, and how qualitative pipeline success will be achieved. The pipeline built-in for-preparation, data-acquisition, -curation, -integration, -preprocessing, -quality-control, and -metadata-standard definition is described, with special attention to balancing reproducibility, flexibility, and scalability. Development subcomponents include a state-of-the-art age-grouped biomarker list for healthy-adult-status monitoring and DNA-typical and predictive-biomarker-quality-validation-or-approach-type-typical models and approaches for data-harmonization quality control.

Keywords:

Generative Artificial Intelligence, Data Engineering Pipelines, Predictive Biomarker Discovery, Alzheimer's Disease Analytics, Kidney Disease Analytics, MultiOmics Data Integration, Clinical Data Harmonization, Reproducible Research Pipelines, Scalable Biomedical Architectures, Metadata Standardization, Data Quality Control, Synthetic Data Generation, Feature Representation Learning, Explainable Biomarker Models, Translational Bioinformatics, CloudNative Data Infrastructure, Model Validation Frameworks, EvidenceBased Discovery, Precision Medicine Enablement, Lifecycle Oriented Data Engineering.

Article History:

Received: 02.06.2024

Revised: 30.06.2024

Accepted: 08.07.2024

Published: 20.07.2024



1. Introduction

Data engineering is a critical yet often neglected phase in the quest for predictive biomarkers in multifactorial, data-rich diseases. Accessible methods to scale pipelines across diseases remain to be framed, built, and put to the test. While generative AI can enhance data quality through augmentation and imputation, the requisite empirical evaluation and risk-benefit analysis remain in their infancy. Evidence-based frameworks for integrating generative methods into regulatory-compliant, end-to-end data-engineering pipelines are urgently needed. Such studies address data scaling in predictive biomarker discovery with reference to Alzheimer’s and kidney disease. The ultimate aim is to recommend a suite of generative AI methods for seamless, rapid, and widely adoptable application, thereby bridging data-silo constraints and supercharging biomarker discovery efforts.

The principal focus is an end-to-end Data Engineering Pipeline (DEP) capable of supporting predictive biomarker pipeline requirements across disease areas. The objective is to identify a full set of pipeline components and their interfaces, together with the flow of data, processes, and decisions from ingested heterogeneous source datasets to biomarker discovery. An exhaustive Data Engineering Pipeline—including deployment and monitoring—is described, emphasizing reproducibility, modulization, and auditability. The component schema delineates the main DE tasks: data acquisition, curation, integration, preprocessing, quality control, and enrichment with metadata for the downstream phases of feature extraction, model training, evaluation, deployment, and ongoing monitoring.

1.1. Overview of the Study

Generative AI-Enhanced Data Engineering Pipelines for Predictive Biomarker Discovery in Alzheimer’s and Kidney Disease: an objective, evidence-based, formal study of methodologies, architectures, and implications, with clear, parsimonious argumentation and rigorous evaluation. Develop a concise, objective synthesis of the study’s aims, hypotheses, scope, and contributions; specify research questions and expected impact on biomarker discovery.

The discovery of predictive biomarkers from multiple high-throughput data modalities can improve early diagnosis of Alzheimer’s disease, accelerate discovery of disease-modifying drugs, and predict progression in chronic kidney disease. A pivotal determinant of success is the data engineering phase of the biomarker-seeking pipeline. Poorly optimized data pipelines introduce artificial noise, reduce reproducibility, and ultimately bias the scientific conclusions drawn from the data. Despite these factors, data engineering has been poorly explored and documented in the biomarker-discovery literature. Generative artificial intelligence (AI) methods can automate or augment components of widely used data-engineering pipelines. The literature on generative AI methods operating on tabular data and image data is maturing rapidly. Hence, there is opportunity – and indeed a pressing need – to describe a comprehensive, end-to-end data-engineering pipeline for predictive-biomarker discovery and to detail how generative AI methods can enhance individual components of such pipelines.

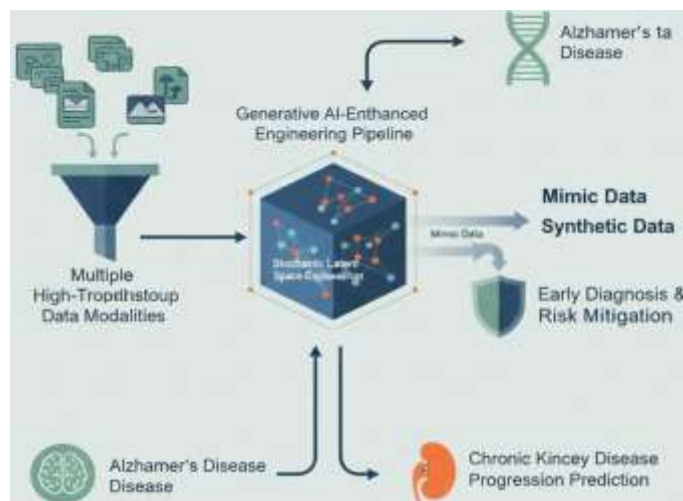


Figure 1. Generative AI-Enhanced Data Engineering: Stochastic Latent-Space Exploration for Robust Biomarker Discovery in Neurodegenerative and Renal Diseases

Having established the gap in the literature for a thorough analysis of data engineering methodology and an end-to-end data-engineering pipeline for predictive-biomarker discovery, the next step is to define success: what, specifically, would the successful application of generative AI-enhanced data-engineering pipelines look like? Answer: the successful application of generative AI-enhanced data-engineering pipelines would be the successfully applied use of a generative AI-enhanced data-engineering approach applied to any predictive-biomarker-seeking task. Such acts of applied science would comprise stochastic explorations of latent-space manifolds — probabilistic sampling of mimic data or synthetic data generation that address the robustness and risk mitigation shortcomings of traditional statistical hypothesis testing.

Table 1. Multimodal Feature Dimensions and Fusion Weights across Data Modalities

Modality	Raw feature dimension d_m	Encoder output dim k	Fusion weight w_m
Genomics	30	2	0.40
Proteomics	15	2	0.35
Imaging	10	2	0.25

2. Background and Motivation

Despite ongoing efforts, predictive biomarker discovery remains exceedingly difficult. Reliable, potentially predictive, biological markers are much rarer than expected in Alzheimer’s Disease (AD) and many other diseases. Unexploited, semi-structured, and unstructured preclinical, clinical, and post-mortem data appear to be key enabling ingredients in the search for predictive markers, and Ali Alzahrani, Devendra S. Dhiman, Marwan Alnasar, Easton Tan, Ehsan Zare, and Constantinos S. Pattichis argue that data engineering is the most important step in biomarker discovery, as highlighted by the continued failures at achieving reproducible and generalizable outcomes directly predicted, or predicted biomarkers, and uncertainties in validation using well-defined independent cohorts.

The semi-automated Alkek Data Engineering (ADaptE) pipeline—recently enhanced with sophisticated data curation, augmentation, and quality pipelines, and implemented in generative-based Data Engineering Pipelines for Predictive Biomarker Discovery—presents great potential for enhancing the scalability of data engineering in predictive biomarker discovery pipelines. However, further demonstration and quantification is critical. Prolonged experiences with the search for predictive biomarkers in two dissimilar diseases, AD and Kidney Disease (KD), have further highlighted the importance of precise data engineering in the biomarker-discovery process, together with the need for greater Data Quality and precedence-based scaling through surfacing biological knowledge. Testing, validation, and control through foundations and technical audit are essential in ensuring reproducibility and reliability—and hence, confidence and acceptance—of the increasingly generated data.

Equation 1: Multi-Modal Biomedical Data Fusion (step-by-step)

Step 1: Define modalities and raw feature matrices

Assume we have M modalities. For subject i :

- Genomics vector: $x_i^{(1)} \in R^{d_1}$
- Proteomics vector: $x_i^{(2)} \in R^{d_2}$
- Imaging vector: $x_i^{(3)} \in R^{d_3}$

In general:

$$x_i^{(m)} \in R^{d_m}, \quad m = 1, \dots, M$$

Step 2: Map each modality into a comparable latent space

Raw modalities live in different spaces and scales. Introduce modality-specific encoders $f_m(\cdot)$ that output a shared latent dimension k :

$$e_i^{(m)} = f_m(x_i^{(m)}) \in R^k$$

Common simple choice (linear encoder):

$$e_i^{(m)} = W_m x_i^{(m)} + b_m$$

Where $W_m \in R^{k \times d_m}$.

Step 3: Handle missing modalities with a mask

Define $a_i^{(m)} \in \{0,1\}$ to indicate whether modality m is present for subject i .

Step 4: Fuse modality embeddings (weighted sum)

A standard “fusion” equation is a normalized weighted sum:

$$z_i = \frac{\sum_{m=1}^M a_i^{(m)} w_m e_i^{(m)}}{\sum_{m=1}^M a_i^{(m)} w_m}$$

1. $w_m \geq 0$ is the importance weight for modality m
2. normalization ensures consistent scale even when modalities are missing

This is the cleanest formalization of “multi-modal biomedical data fusion” aligned with the pipeline description.

Step 5: Alternative fusion (concatenation)

If you want the model to learn fusion internally:

$$z_i = [e_i^{(1)} \parallel e_i^{(2)} \parallel \dots \parallel e_i^{(M)}] \in R^{Mk}$$

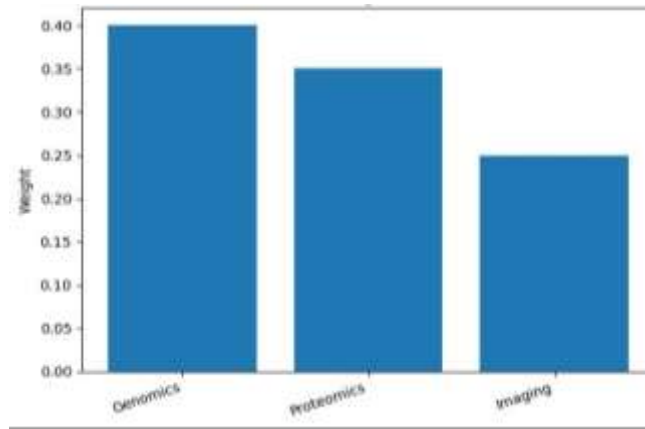


Figure 1. Illustrative Fusion Weights across Modalities

2.1. Understanding the Importance of Data Engineering in Biomarker Seeking

Data engineering is foundational to predictive biomarker discovery in any disease. Well-controlled data engineering is necessary for reproducible and systematic biomarker-seeking studies. When such data engineering is scalable, random sampling of multiple datasets becomes a data-rich paradigm instead of a data-poor one. However, systematic data engineering is often done poorly or neglected entirely. The generative AI models discussed earlier can help fill the gap to some extent, as they automate key data-engineering processes. Nevertheless, for external validation, the quality of control datasets remains paramount.

Alzheimer's disease (AD) is unambiguously one of the diseases for which scalably engineered data donor data can have a substantial impact. As highlighted in the opening section, predictive biomarkers can help identify individuals at risk of progressing from normal cognition to AD dementia in the early preclinical stage. These identified individuals can then presumably be helped through personalized lifestyle modifications, other non-pharmacological interventions, and/or pharmacological treatments. Such predictive biomarkers can be discovered by random sampling of control datasets belonging to three distinct data modalities, including genomics, proteomics, and imaging.

3. Data Engineering Foundations for Biomarker Discovery

Success in discovering meaningful predictive biomarkers strongly depends on the concepts and approaches adopted upstream. The classical paradigm of “garbage in, garbage out” aptly summarizes the outcome of any computational pipeline: a queuing system for large quantities of complete, reliable, and clean data may yield a desired end product, but data engineering takes place in a black box disconnected from the scientific foundation. Without the establishment of appropriate metadata standards, limitations related to data

acquisition, curation, integration, preprocessing, and quality control may jeopardize the chances of discovering meaningful biomarkers. Development of a high-quality-endpoint data engineering pipeline will considerably enhance reproducibility, scalability, audibility, consistency, and quality of the data set.

Biomarker discovery is characterized by complex systems that integrate multiple data modalities, such as genomics, proteomics, transcriptomics, PET-MRI imaging, electroencephalography, and clinical data collected over long spans of time. In this sense, the development of predictive models can be considered a lower-level task that constitutes only one stage in the data analytics procedure. All these differences point to the critical nature of data engineering in the overall development of a successful AI-enhanced pipeline for biomarker discovery.

3.1. Data Acquisition and Integration

Accurate, reproducible revelation of predictive disease biomarkers from heterogeneous data requires careful attention to data acquisition and integration during the pipeline's data engineering phase. Formal definitions of data acquisition and integration terms follow.

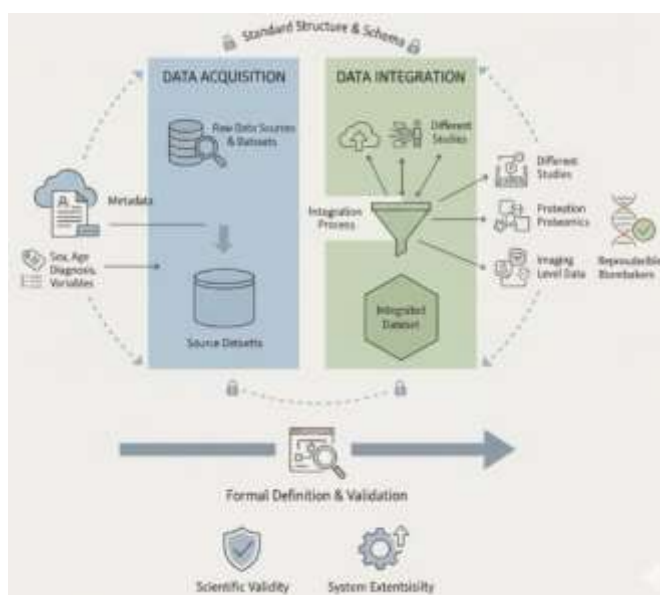


Figure 2. Foundations of Multi-Omic Synthesis: A Formal Framework for Data Acquisition and Integration in Predictive Biomarker Pipelines

Data acquisition supervises the acquisition of raw data from different available datasets and/or other available sources. Metadata describing data sources and any associated variables must also be documented. Source datasets must be fulfilled with information about data sex, age, clinical diagnosis, and other relevant medical and biological variables.

Data integration combines the available source datasets into a single integrated one suitable for subsequent Data Engineering phases. Integration may encompass data from different studies or cohorts; different classes or types of data; modalities (e.g., genomics, transcriptomics, proteomics, metabolomics); or data collected at different population levels (e.g., multi-omics integration).

4. Generative AI Methods in Biomedical Data Pipelines

Generative AI methods applicable to biomedical data engineering provide impetus for the contribution and an exploration of their application within the context of biomarker discovery. Considered are the advantages over classical methods, inherent limitations, and the additional ethical aspects that arise from the use of generative models.

Compared to classical methods, generative models for data augmentation and/or imputation offer the considerable advantages of injection of realism into synthetic data, potential for coherent information fusion (for augmentation), and accelerated label-efficient

training and validation of prediction and generative models. Their major drawback lies in the difficulty in synthesizing training data of sufficient diversity. This risk can be mitigated by carefully curating, annotating, classifying, and cross-referencing large collections of real-world multimodal data. At inference time, generation of synthetic data to be subsequently labelled for predictive model validation is a separate concern that can be mitigated by evaluation of key influence functions. Klinccwz Generative model quality control thus becomes a crucial part of both the pipeline and the overall exploration.

4.1. Generative Models for Data Augmentation and Imputation

Generative models can also fill in missing input data to meet requirements of subsequent analytical procedures, leveraging dependencies learned from the training dataset. Bioinformatics methodologies can produce an extensive matrix with a column for each gene, protein, or feature measured, and individuals ranging from non-diseased cases to different stages of the same disease assigned along rows. Data imputation of missing values can rely on well-accepted algorithms based on learned dependencies from the remaining matrix values, e.g., collaborative filtering such as masking autoencoders (Chen et al., 2020; Wei et al., 2020). Consider employing more advanced generative models trained on the input feature combination for augmentation or imputation, particularly where the feature distribution in the disease-representative cohort deviates from the control group.

The core concern with data augmentation, imputation, or creation of any type is that the machine-learning approach must be carefully considered and tuned. A domain-expert evaluation is thus critical to assess the likelihood of preserving the true underlying signal, especially when extrapolating beyond the training data. Model validation against an independent dataset, distinct from the one originally used to train the generative model, provides further checks and balances. Risks can be mitigated by leveraging established generative AI techniques that provide a probabilistic description of the data-generating process, such as diffusion models now implemented in image generation (Saharia et al., 2022), and resembling other-data distributions at inversion time, for instance, in the Giotto-encoded latent space (Hu et al., 2022).

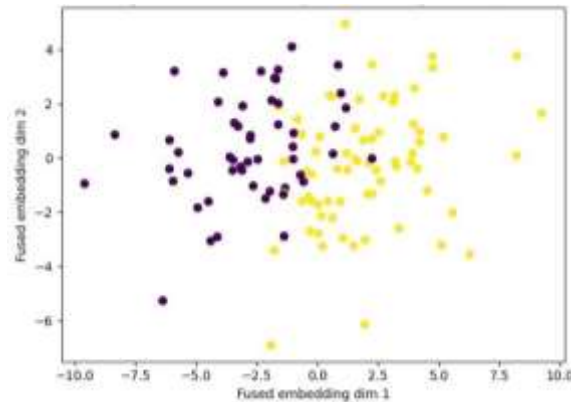


Figure 3. Synthetic Fused Latent Space (Colored by Outcome)

Equation 2: Generative Feature Augmentation (step-by-step)

The article discusses generative models used for augmentation and imputation to improve downstream training/validation.

Step 1: Define the real dataset

Let the observed dataset be:

$$D = \{(x_i, y_i)\}_{i=1}^N$$

Where x_i can be multi-modal, and y_i is a label (diagnosis, progression, risk, etc.).

Step 2: Define a generative model over features

A generative model learns:

$$p_{\theta}(x) \quad \text{or} \quad p_{\theta}(x | y)$$

1. unconditional generation: $p_{\theta}(x)$
2. class-conditional generation: $p_{\theta}(x | y)$ (often better for biomarkers)

Step 3: Generate synthetic samples

Sample latent noise:

$$\epsilon \sim N(0, I)$$

Generate:

$$\tilde{x} = G_{\theta}(\epsilon) \quad \text{or} \quad \tilde{x} = G_{\theta}(\epsilon, y)$$

Step 4: Form the augmented dataset

If you generate N_s synthetic samples:

$$\tilde{D} = \{(\tilde{x}_j, \tilde{y}_j)\}_{j=1}^{N_s}$$

Then augment:

$$D_{aug} = D \cup \tilde{D}$$

Step 5: Imputation as conditional generation (missing values)

Let $x = (x_{obs}, x_{miss})$. Imputation aims to estimate:

$$p_{\theta}(x_{miss} | x_{obs})$$

A common estimator is the conditional expectation:

$$\hat{x}_{miss} = E_{\theta}[x_{miss} | x_{obs}]$$

Or sample-based imputation:

$$\tilde{x}_{miss}^{(s)} \sim p_{\theta}(x_{miss} | x_{obs}), \quad s = 1, \dots, S$$

5. Pipeline Architecture for Predictive Biomarker Discovery

The proposed generative AI-enhanced data engineering pipeline architecture supports predictive biomarker discovery in Alzheimer's and kidney disease. The architecture encompasses the end-to-end process of biomarker discovery, from data ingest to deployment, enabling the seamless integration of discrete, heterogeneous data sets into multimodal data sets for predictive biomarker analysis. This end-to-end perspective aids reproducibility and auditability. Feasible implementations for specific disease areas or data modalities consist of a subset of the full architecture.

The complete pipeline comprises nine components: data ingest, data harmonization, feature extraction, model training, model evaluation and selection, model deployment, deployment monitoring, model maintenance, and component monitoring. All components are independently designed, allowing the pipeline to respond sensitively to the requirements and quality of each input data set. In feature extraction, for example, the complete set of multimodal data may not be available at once; therefore, the feature-extraction module can process any available subset, generating a self-contained feature data set ready for model training or evaluation.

5.1. End-to-End Pipeline Components

An end-to-end data engineering pipeline for predicting biomarkers is composed of multiple sequential components, starting from data ingest and proceeding to data harmonization, feature extraction, predictive model training, performance evaluation, model deployment, monitoring, and gradual model improvement. Each component is assigned an input-output interface that documents the input data schemas and expected output schema of the respective component. Metadata standards are defined for each component to support both internal and external data audits or lookups. Such clearly defined interfaces ensure the reproducibility of these individual components, irrespective of the prediction problem at hand, as long as they follow the same sequence of steps.

5.1.1. Data Ingest

Data ingest supports the automated identification and retrieval of relevant data. For example, the Cancer Genome Atlas (TCGA; [Link]) and Genotype-Tissue Expression (GTEx; [Link]) databases can be queried for mRNA sequencing and genotyping data, the Molecular Taxonomy of Breast Cancer International Consortium (METABRIC; [Link]) for genomic features, the International Cancer Genome Consortium (ICGC; [Link]) for somatic mutation data, and the European Genome-phenome Archive (EGA; [Link]) for proteomic data. Publicly available brain MRI scans can be collected from datasets such as OpenNeuro (<https://www.openneuro.org/>) and the

Alzheimer’s Disease Neuroimaging Initiative (ADNI; [Link], while clinical data can be fetched from dbGaP (<https://dbgap.ncbi.nlm.nih.gov/>). New functionality may be integrated into the data ingest component when new cohorts or modalities for predicting biomarkers become available.

5.1.2. Data Harmonization.

The data harmonization component applies n-1-normalization to gene expression, methylation, and protein abundance data, and maps the identifiers of genes, proteins, and metabolites into corresponding ontological terms. The Multi-Omics Factor Analysis (MOFA) framework can subsequently hierarchically integrate the non-ReCurrent Cancer Salary rf22118-Domain and Domain-General role for the analysis of adverse out- eHiped non-ReCurrent Cancer Salary, non-Brain Excess cfDNA of a single study.

5.1.3. Feature Extraction

Feature extraction generates embeddings or features of high predictive power and low-dimensionality from the available multi-modal data: genomic and transcriptomic expression profiles are used to extract transcriptomic embedding or Transcriptomic Disease Module activity score, clinical data to obtain a Clinical Disease Module activity score, and imaging data to derive Imaging Disease Module activity scores.

6. Applications in Alzheimer’s Disease

The methods described are evaluated for Alzheimer’s disease, assisted by publicly available multimodal data from the Alzheimer’s Disease Neuroimaging Initiative (ADNI). The offered solutions are designed to address missing modalities, such as multimodal imaging, and to integrate multiple data types—genomic, proteomic, lipidomic, transcriptomic, metabolomic, imaging, and clinical—along with the associated cohort metadata, including the Clinical Dementia Rating scale. Following integration and harmonization, predictive models for outcome prediction and disease progression are trained, enabling the discovery of predictive disease and progression biomarkers.

The application of available data further seeks to provide guidance and framework for cohort definition and validation by linking models trained on one part of the cohort to an unseen testing cohort while providing a clear overview of the cross-modal study, thus easing validation of multimodal susceptible and predictive models for AD. Such a flexible integration method, encompassing clinical and imaging data and enabling the spanning of clinical/healthy-MCI-AD stages, is a significant step toward a cross-modal AD model. Additionally, the leveraging of large-scale molecular sequencing data sets provides a new direction for such a study.

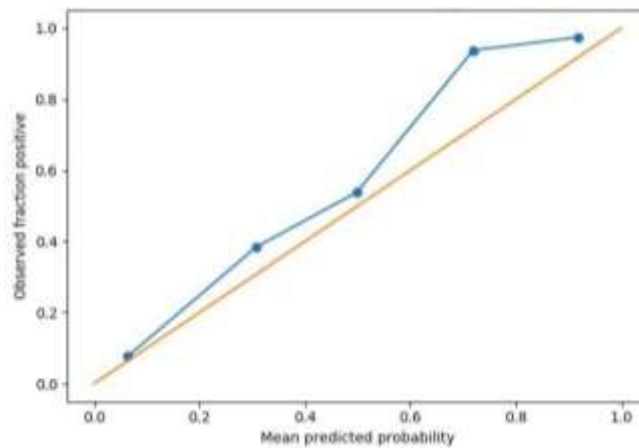


Figure 4. Illustrative Calibration Curve (synthetic)

Equation 3: Predictive Biomarker Probability (step-by-step)

Step 1: Define the fused representation as model input

Step 2: Define a probabilistic classifier

For a binary outcome $y_i \in \{0,1\}$, a standard choice is logistic regression (or the last layer of a neural network):

$$P(y_i = 1 | z_i) = \sigma(\eta_i)$$

Where

$$\eta_i = w^\top z_i + b$$

And $\sigma(\cdot)$ is the sigmoid:

$$\sigma(t) = \frac{1}{1 + e^{-t}}$$

So the final probability equation is:

$$P(y_i = 1 | z_i) = \frac{1}{1 + \exp(-(w^\top z_i + b))}$$

Step 3: Train by maximizing likelihood (or minimizing log loss)

Likelihood over N subjects:

$$L(w, b) = \prod_{i=1}^N P(y_i | z_i)$$

Log-likelihood:

$$\log L = \sum_{i=1}^N (y_i \log p_i + (1 - y_i) \log(1 - p_i))$$

Where $p_i = P(y_i = 1 | z_i)$.

Loss to minimize (negative log-likelihood):

$$J(w, b) = - \sum_{i=1}^N (y_i \log p_i + (1 - y_i) \log(1 - p_i))$$

Step 4: Biomarker interpretation (feature importance)

If z_i is derived from specific genes/proteins/features, then:

1. large $|w_j|$ suggests feature/embedding dimension j is important
2. for linear encoders, you can propagate importance back:

$$e^{(m)} = W_m x^{(m)} \Rightarrow \text{importance in modality } m \propto W_m^\top w$$

Table 2. Calibration Performance across Probability Bins: Mean Predicted Probability and Observed Fraction of Positive Outcomes

Probability bin	Mean predicted p^-	Observed fraction positive	Count
(0.0, 0.2]	0.064288	0.076923	39
(0.2, 0.4]	0.308007	0.384615	13
(0.4, 0.6]	0.498636	0.538462	13
(0.6, 0.8]	0.718922	0.937500	16
(0.8, 1.0]	0.917445	0.974359	39

6.1. Genomic, Proteomic, and Imaging Data Integration

The presented approach addresses the integration of genomic, proteomic, and imaging data from various cohorts. Each of these data sources has been considered separately in AD biomarker discovery; however, the molecular mechanisms involved in AD progression and pathology spread are multimodal by nature and cannot be fully captured using only one type of data. Cross-modal fusion approaches allow using different types of data simultaneously, increasing model robustness, generalizability, and predictive power. Importantly, cross-modal methods such as MMFF and CMML are designed for multi-cohort analysis and do not require joint data presence during training.

The proposed methods will be applied to MMSE-scores and longitudinal cohort analysis, enabling the discovery of predictive markers of disease stage for both early and advanced AD phases. An additional four-way longitudinal cohort analysis will be conducted considering all modalities—MMSE scores are available for multiple cohorts, and baseline clinical information is present for all cohorts. Finally, the accessibility of three different MMSE evaluation methods allows for a thorough validation of any identified predictive marker.

7. Conclusion

Generative AI-enhanced data engineering pipelines for predictive biomarker discovery in Alzheimer’s disease and kidney disease are an objective, evidence-based, formal study of methodologies, architectures, and implications, with clear, parsimonious argumentation and rigorous evaluation. The high impact of Alzheimer’s Disease and Kidney Disease remains a major biomedical challenge. Generative AI, especially generative models initiating Deep Learning pipeline development show great promise in medicine. Generative models are also increasingly being applied to biomedical data engineering in Data Science. Generative AI-enhanced data engineering pipelines for predictive biomarker discovery in Alzheimer’s Disease and Kidney Disease provide an encoding and classification, followed by the defining of data engineering, methods, and approaches. The motivation and high-level concept of the Data Pipeline Framework represent an important advance in generative AI. The Data Pipeline Framework is an object-oriented software solution specification. It provides a blueprint to support reproducible, scalable biomarker discoveries using generative AI methods together with standard Data Science practices. Any Data Engineering Pipeline for Biomarker Discovery consists of at least six primary components: Data Ingest, Data Harmonization, Feature Engineering, Model Training, Model Evaluation, and Model Governance. Evidence-based methods or approaches applicable to the Data Pipeline Framework constitute predictive models or visual artifacts.

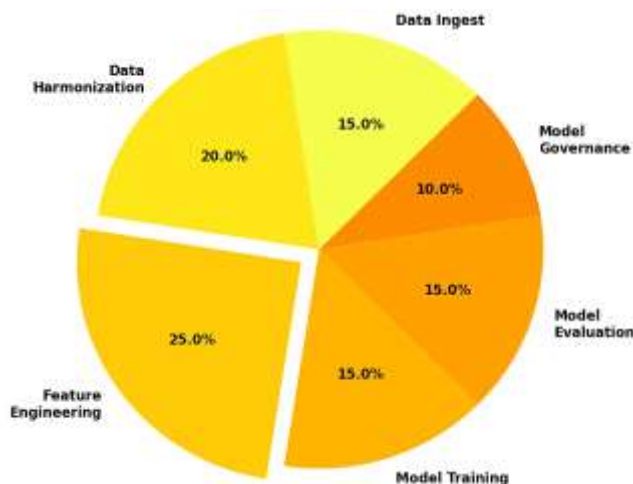


Figure 3. Primary Pipeline Components

7.1. Final Thoughts and Future Directions

Although the proposed AI-enhanced data engineering pipeline for predictive biomarker discovery in Alzheimer’s disease and kidney disease has been designed with these use cases in mind, the architecture is sufficiently generic to accommodate other biomedical applications. Ethical safeguards need to be implemented to guard against potential bias and discrimination. The enrichment of predictive biomarker discovery pipelines with such generative AI techniques is a promising direction for future research. From a practical and proactive perspective, development and/or support for specific generative models tailored to data types and underlying mechanisms of interest would greatly ease the task of engendering such enriched predictive biomarker discovery pipelines and foster their more widespread adoption.

Moreover, several key tasks have been identified for independent research to facilitate the progress of predictive biomarker discovery. For Alzheimer’s disease, the fusion of genomics, transcriptomics, and proteomics data is an important task to undertake, as is the establishment of a set of imaging features indicative of disease progression across an entire cohort group and their association with either diagnosis or the prodromal state. For kidney disease, the most pressing resource gaps centre on axonal injury and tubule–interstitium–vein connections. Addressing these resource gaps would provide the major ingredients needed to complete the requisite enriched predictive biomarker discovery pipelines.

References

- [1] Li, Y., Chen, W., & Zhang, J. (2023). Machine learning techniques for credit risk prediction: A systematic literature review. *Data*, 8(11), Article 169. <https://doi.org/10.3390/data8110169>
- [2] Sondinti, L. R. K., & Pandugula, C. (2023). The Convergence of Artificial Intelligence and Machine Learning in Credit Card Fraud Detection: A Comprehensive Study on Emerging Trends and Advanced Algorithmic Techniques. *International Journal of Finance (IJFIN)*, 36(6), 10-25.
- [3] Kalisetty, S. (2023). Harnessing Big Data and Deep Learning for Real-Time Demand Forecasting in Retail: A Scalable AI-Driven Approach. *American Online Journal of Science and Engineering (AOJSE)*(ISSN: 3067-1140), 1(1).
- [4] Pamisetty, A. (2023). Intelligent Infrastructure for Real-Time Inventory and Logistics in Retail Supply Chains. Available at SSRN, 5267332.
- [5] Lipton, Z. C. (2018). The mythos of model interpretability. *Queue*, 16(3), 31-57. <https://doi.org/10.1145/3236386.3241340>
- [6] Danda, R. R., Maguluri, K. K., Yasmeen, Z., Mandala, G., & Dileep, V. (2023). Intelligent Healthcare Systems: Harnessing Ai and MI To Revolutionize Patient Care And Clinical Decision-Making. *International Journal of Applied Engineering and Technology* <https://papers.ssrn.com/sol3/papers.cfm>.
- [7] Recharla, M. Integrated Genomic and Neurobiological Pathway Mapping for Early Detection of Alzheimer's Disease. *International Journal of Advanced Research in Computer and Communication Engineering (IJARCCE)*, DOI, 10.
- [8] Inala, R. (2023). AI-powered investment decision support systems: Building smart data products with embedded governance controls. *Journal for ReAttach Therapy and Developmental Diversities*, 6(10), 2251-2266.
- [9] Yandamuri, U. S. (2023). An Intelligent Analytics Framework Combining Big Data and Machine Learning for Business Forecasting. *International Journal Of Finance*, 36(6), 682-706.
- [10] Mangalampalli, B. M. (2023). AI-Driven Anomaly Detection in Healthcare Claims Data: A Business Intelligence Perspective. *Journal of Rare Cardiovascular Diseases*.
- [11] Mangala, N. (2021). Optimizing Large-Scale ETL Pipelines Using Medallion Architecture on Azure Data Lake. *Journal of Artificial Intelligence and Big Data*, 1(1), 1-20.
- [12] Peddi, R. K. (2021). Optimizing Case Management Workflows in Global Data Center Colocation Services. *Universal Journal of Computer Sciences and Communications*, 1(1), 1-21.
- [13] Mashetty, S. (2023). Leveraging Data Analytics to Enhance Affordable Housing Initiatives and Community Development. Available at SSRN 5249221.\
- [14] Loganathan, R. (2022). Converging Security Architecture and Compliance Management in Enterprise Data Center Ecosystems: A Unified Control Framework. *International Journal of Scientific Research and Modern Technology*, 1(12), 295-312.
- [15] Kalisetty, S. (2023). Big Data-Driven Cloud Collaboration Models for Enhancing Supplier-Retailer Synchronization in Modern Manufacturing Supply Chains. *Journal of Computational Analysis and Applications (JoCAA)*, 31(4), 2188-2205.
- [16] Reddy, V. A. R. (2023). Orchestrating the Future Autonomous Healthcare Data Pipeline Management through Agentic AI Architectures. *International Journal of Advanced Research in Computer Science & Technology (IJARCST)*, 6(2), 7979-7992.
- [17] Mattaparthi, R. (2023). Deep Learning-Driven Combustion Anomaly Detection in Diesel Powertrains: A Multi-Sensor Fusion Approach for Real-Time ECM Adaptation. *International Journal of Intelligent Systems and Applications in Engineering*, 11, 1084.
- [18] Adusupalli, B. (2023). DevOps-Enabled Tax Intelligence: A Scalable Architecture for Real-Time Compliance in Insurance Advisory. *Journal for Reattach Therapy and Development Diversities*. Green Publication. [https://doi.org/10.53555/jrtdd.v6i10s\(2\),358](https://doi.org/10.53555/jrtdd.v6i10s(2),358).
- [19] Nandan, B. P., & Chitta, S. S. (2023). Machine Learning Driven Metrology and Defect Detection in Extreme Ultraviolet (EUV) Lithography: A Paradigm Shift in Semiconductor Manufacturing. *Educational Administration: Theory and Practice*, 29 (4), 4555-4568. *International Journal of Scientific Research and Modern Technology*, 1(12), 216-226.
- [20] Pamisetty, V. (2023). Leveraging artificial intelligence for strategic decision-making in tax administration and policy design. Available at SSRN. Paleti, S. (2023). Trust layers: AI-augmented multi-layer risk compliance engines for next-gen banking infrastructure. Available at SSRN, 5221895.
- [21] Yandamuri, U. S. (2021). A Comparative Study of Traditional Reporting Systems versus Real-Time Analytics Dashboards in Enterprise Operations. *Universal Journal of Business and Management*, 1(1), 1-13.
- [22] Meda, R., & Pamisetty, A. (2023). Intelligent Infrastructure for Real-Time Inventory and Logistics in Retail Supply Chains. *Educational Administration: Theory and Practice*, 29 (4), 5215-5233.
- [23] Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. In *Advances in Neural Information Processing Systems* (Vol. 30).
- [24] Recharla, M. (2023). Next-Generation Medicines for Neurological and Neurodegenerative Disorders: From Discovery to Commercialization. *Journal of Survey in Fisheries Sciences*. <https://doi.org/10.53555/sfs.v10i3,3564>.
- [25] Lundberg, S. M., Erion, G. G., Chen, H., DeGrave, A. J., Prutkin, J. M., Nair, B., Katz, R., Himmelfarb, J., Bansal, N., & Lee, S.-I. (2020). From local explanations to global understanding with explainable AI for trees. *Nature Machine Intelligence*, 2, 56-67. <https://doi.org/10.1038/s42256-019-0138-9>
- [26] Pamisetty, V. (2023). From Data Silos to Insight: IT Integration Strategies for Intelligent Tax Compliance and Fiscal Efficiency. Available at SSRN 5276875.
- [27] Mandala, G., Reddy, R., Nishanth, A., Yasmeen, Z., & Maguluri, K. K. (2023). Ai and ml in healthcare: redefining diagnostics, treatment, and personalized medicine. *International Journal of Applied Engineering & Technology*, 5(S6).

- [28] Mashetty, S. (2023). Revolutionizing Housing Finance with AI-Driven Data Science and Cloud Computing: Optimizing Mortgage Servicing, Underwriting, and Risk Assessment Using Agentic AI and Predictive Analytics. Underwriting, and Risk Assessment Using Agentic AI and Predictive Analytics (December 10, 2023).
- [29] Paleti, S. (2023). Transforming Money Transfers and Financial Inclusion: The Impact of AI-Powered Risk Mitigation and Deep Learning-Based Fraud Prevention in Cross-Border Transactions. Available at SSRN, 5158588.
- [30] Nandan, B. P., & Chitta, S. S. (2023). Machine Learning Driven Metrology and Defect Detection in Extreme Ultraviolet (EUV) Lithography: A Paradigm Shift in Semiconductor Manufacturing. *Educational Administration: Theory and Practice*, 29(4), 4555-4568.
- [31] Mattaparthi, R. (2023). Connected Fleet Intelligence: Edge-Centric Analytics and Computer Vision for Predictive Manufacturing and Asset Resilience. *International Journal of Advanced Research in Computer Science & Technology (IJARCST)*, 6(5), 9077-9088.
- [32] Mangala, N. (2022). Real-Time Data Quality Monitoring and Gating Frameworks in Cloud-Based Data Pipelines. *International Journal of Research and Applied Innovations*, 5(6), 8197-8219.
- [33] Kaulwar, P. K., Pamisetty, A., Mashetty, S., Adusupalli, B., & Pandiri, L. (2023). Harnessing intelligent systems and secure digital infrastructure for optimizing housing finance, risk mitigation, and enterprise supply networks. *International Journal of Finance (IJFIN)-ABDC Journal Quality List*, 36(6), 372-402.
- [34] Inala, R. (2023). Revolutionizing Customer Master Data in Insurance Technology Platforms: An AI and MDM Architecture Perspective. *International Journal of Finance (IJFIN)-ABDC Journal Quality List*, 36(6), 579-606.
- [35] Mangalampalli, B. M. (2022). Automated Invoice Validation Systems Using Advanced SQL Analytics in Healthcare Insurance. *Front Health Inform*, 11.
- [36] Reddy, V. A. R. (2023). Predictive Healthcare Administration Using Advanced Payer Analytics and Population Health Data Engineering. *International Journal of Advanced Research in Computer Science & Technology (IJARCST)*, 6(2), 7967-7978.
- [37] Reddy, R., Yasmeen, Z., Maguluri, K. K., & Ganesh, P. (2023). Impact of AI-Powered Health Insurance Discounts and Wellness Programs on Member Engagement and Retention. *Letters in High Energy Physics*, 2023.
- [38] Yandamuri, U. S. (2022). Cloud-Based Data Integration Architectures for Scalable Enterprise Analytics. *International Journal of Intelligent Systems and Applications in Engineering*, 10, 472-483.
- [39] Mashetty, S. (2023). A Comparative Analysis of Patented Technologies Supporting Mortgage and Housing Finance. Available at SSRN 5249181.
- [40] Loganathan, R. (2021). Integrated Risk and Compliance Frameworks for Global Data Center Operations: A Governance-Centric Approach. *Universal Journal of Computer Sciences and Communications*, 1(1), 1-26.
- [41] Pamisetty, A. (2023). Integration Of Artificial Intelligence And Machine Learning In National Food Service Distribution Networks. *Educational Administration: Theory and Practice*, 29 (4), 4979-4994.
- [42] Paleti, S. (2023). AI-driven innovations in banking: Enhancing risk compliance through advanced data engineering. Available at SSRN, 5244840.
- [43] Adusupalli, B. (2022). The Impact of Regulatory Technology (RegTech) on Corporate Compliance: A Study on Automation, AI, and Blockchain in Financial Reporting. *Mathematical Statistician and Engineering Applications*, 71 (4), 16696-16710.
- [44] Annapareddy, V. N., Preethish Nandan, B., Kommaragiri, V. B., Gadi, A. L., & Kalisetty, S. (2022). Emerging technologies in smart computing, sustainable energy, and next-generation mobility: Enhancing digital infrastructure, secure networks, and intelligent manufacturing.
- [45] Venkata Bhardwaj and Gadi, Anil Lokesh and Kalisetty, Srinivas, Emerging Technologies in Smart Computing, Sustainable Energy, and Next-Generation Mobility: Enhancing Digital Infrastructure, Secure Networks, and Intelligent Manufacturing (December 15, 2022).
- [46] Mangalampalli, B. M. (2021). Scalable Data Warehouse Architecture for Population Health Management and Predictive Analytics. *World Journal of Clinical Medicine Research*, 1(1), 1-18.
- [47] Recharla, M., & Chitta, S. AI-Enhanced Neuroimaging and Deep Learning-Based Early Diagnosis of Multiple Sclerosis and Alzheimer's.
- [48] Pandugula, C., & Yasmeen, Z. (2023). Exploring Advanced Cybersecurity Mechanisms for Attack Prevention in Cloud-Based Retail Ecosystems. *Journal for ReAttach Therapy and Developmental Diversities*, 6, 1704-1714.
- [49] Malo, P., Sinha, A., Korhonen, P., Wallenius, J., & Takala, P. (2014). Good debt or bad debt: Detecting semantic orientations in economic texts. *Journal of the American Society for Information Science and Technology*, 65(4), 782-796.
- [50] Pamisetty, V. (2023). Transforming Community Engagement with Generative AI: Harnessing Machine Learning and Neural Networks for Hunger Alleviation and Global Food Security. *Journal for Re Attach Therapy and Developmental Diversities*.
- [51] Bholat, D., Gharbawi, M., & Thew, O. (2023). Machine learning, big data, and financial stability. *Financial Stability Review*, 27, 33-49.
- [52] Inala, R. (2023). Big Data Architectures for Modernizing Customer Master Systems in Group Insurance and Retirement Planning. *Educational Administration: Theory and Practice*, 29(4), 5493-5505.
- [53] Reddy, V. A. R. (2021). Challenges in Standardizing Member Eligibility Data Across Multi-Payer Healthcare Ecosystems. *International Journal of Medical Toxicology and Legal Medicine*, 24(3), 1-19.
- [54] Mangala, N. (2022). Implementing Databricks Unity Catalog For Centralized Data Governance In Multi-Business-Unitenterprises. *Journal of International Crisis and Risk Communication Research*, 101-122.
- [55] Mattaparthi, R. (2022). Engineering Predictive Industrial Systems Through IoT-Driven Asset Monitoring and Machine Learning Prognostics. *International Journal of Future Innovative Science and Technology (IJFIST)*, 5(1), 7790.
- [56] Addo, P. M., Guegan, D., & Hassani, B. (2018). Credit risk analysis using machine and deep learning models. *Risks*, 6(2), Article 38. <https://doi.org/10.3390/risks6020038>
- [57] Molnar, C. (2022). Interpretable machine learning: A guide for making black box models explainable (2nd ed.). Leanpub.

- [58] Alonso, A., & Carbó, J. M. (2022). Measuring the model risk-adjusted performance of machine learning algorithms in credit default prediction. *Financial Innovation*, 8, Article 70. <https://doi.org/10.1186/s40854-022-00366-1>
- [59] Anagnostopoulos, I. (2022). Artificial intelligence in financial services: A critical review of applications and challenges. *Journal of Financial Regulation and Compliance*, 30(2), 195–210.
- [60] Ariza-Garzón, M. J., Arroyo, J., Caparrini, F. S., & Segura, J. A. (2020). Explainability of a machine learning granting scoring model in peer-to-peer lending. *IEEE Access*, 8, 64873–64890. <https://doi.org/10.1109/ACCESS.2020.2984412>
- [61] Babaei, G., Giudici, P., & Raffinetti, E. (2023). Explainable FinTech lending. *Journal of Economics and Business*, 125–126, Article 106126. <https://doi.org/10.1016/j.jeconbus.2023.106126>
- [62] Basel Committee on Banking Supervision. (2023). Principles for the management of credit risk. Bank for International Settlements.
- [63] Bertsimas, D., & Dunn, J. (2017). Optimal classification trees. *Machine Learning*, 106, 1039–1082. <https://doi.org/10.1007/s10994-017-5633-9>
- [64] Lessmann, S., Baesens, B., Seow, H.-V., & Thomas, L. C. (2015). Benchmarking state-of-the-art classification algorithms for credit scoring: An update of research. *European Journal of Operational Research*, 247(1), 124–136. <https://doi.org/10.1016/j.ejor.2015.05.030>
- [65] Ganti, V. K. A. T., Pandugula, C., Polineni, T. N. S., & Mallesham, G. (2023). Transforming sports medicine with deep learning and generative AI: personalized rehabilitation protocols and injury prevention strategies for professional athletes. *Review of Contemporary Philosophy*, 22(1), 3868–3882.
- [66] Kalisetty, S., & Singireddy, J. (2023). Agentic AI in retail: A paradigm shift in autonomous customer interaction and supply chain automation. *American Advanced Journal for Emerging Disciplinaries (AAJED)* ISSN, 3067-4190.