



Reinforced Risk Prediction With Budget Constraint Using Irregularly Measured Data From Electronic Health Records

Yinghao Pan, Eric B. Laber, Maureen A. Smith & Ying-Qi Zhao

To cite this article: Yinghao Pan, Eric B. Laber, Maureen A. Smith & Ying-Qi Zhao (2023) Reinforced Risk Prediction With Budget Constraint Using Irregularly Measured Data From Electronic Health Records, Journal of the American Statistical Association, 118:542, 1090-1101, DOI: [10.1080/01621459.2021.1978467](https://doi.org/10.1080/01621459.2021.1978467)

To link to this article: <https://doi.org/10.1080/01621459.2021.1978467>

 View supplementary material 

 Published online: 30 Nov 2021.

 Submit your article to this journal 

 Article views: 697

 View related articles 

 View Crossmark data 



Reinforced Risk Prediction With Budget Constraint Using Irregularly Measured Data From Electronic Health Records

Yinghao Pan^a, Eric B. Laber^b, Maureen A. Smith^c, and Ying-Qi Zhao^d

^aDepartment of Mathematics and Statistics, University of North Carolina at Charlotte, Charlotte, NC; ^bDepartment of Statistics, North Carolina State University, Raleigh, NC; ^cDepartments of Population Health Sciences and Family Medicine, University of Wisconsin-Madison, Madison, WI; ^dPublic Health Sciences Division, Fred Hutchinson Cancer Research Center, Seattle, WA

ABSTRACT

Uncontrolled glycated hemoglobin (HbA1c) levels are associated with adverse events among complex diabetic patients. These adverse events present serious health risks to affected patients and are associated with significant financial costs. Thus, a high-quality predictive model that could identify high-risk patients so as to inform preventative treatment has the potential to improve patient outcomes while reducing healthcare costs. Because the biomarker information needed to predict risk is costly and burdensome, it is desirable that such a model collect only as much information as is needed on each patient so as to render an accurate prediction. We propose a sequential predictive model that uses accumulating patient longitudinal data to classify patients as: high-risk, low-risk, or uncertain. Patients classified as high-risk are then recommended to receive preventative treatment and those classified as low-risk are recommended to standard care. Patients classified as uncertain are monitored until a high-risk or low-risk determination is made. We construct the model using claims and enrollment files from Medicare, linked with patient electronic health records (EHR) data. The proposed model uses functional principal components to accommodate noisy longitudinal data and weighting to deal with missingness and sampling bias. The proposed method demonstrates higher predictive accuracy and lower cost than competing methods in a series of simulation experiments and application to data on complex patients with diabetes. Supplementary materials for this article are available online.

ARTICLE HISTORY

Received August 2019
Accepted August 2021

KEYWORDS

Classification with reject option; Cost-effective; Electronic health records; Functional principal component analysis; Reinforcement learning

1. Introduction

Diabetes affects 9.4% of the U.S. population and is associated with costly adverse healthcare outcomes (Centers for Disease Control and Prevention 2017). For patients with diabetes, a controlled glycated hemoglobin (HbA1c) level ($\leq 7\%$) is known to reduce the risk of microvascular complications in both type 1 and type 2 diabetes (ADVANCE Collaborative Group 2008). Thus, there is tremendous clinical interest in identifying patients who are at increased risk of chronic high HbA1c. Traditional risk prediction models use baseline factors known to be associated with disease, for example, the Framingham risk scores for predicting cardiovascular outcomes (Wilson et al. 1998). In the context of diabetes, biomarker information, for example, BMI, blood pressure, metabolic profile, lifestyle factors, etc., have been used to predict cross-sectional or short-term HbA1c control (Park et al. 2002; Chien et al. 2010; Kwon et al. 2014). However, predictive models for long-term control have not been widely studied; this is due in part to the lack of sufficiently rich data to construct and evaluate such models.

The recent curation and linking of large observational data sets has made it possible to study a wide range of chronic diseases in complex patients with diabetes. We use such data to construct a predictive model for long-term HbA1c control in complex patients with diabetes. The data we consider comprises claims and enrollment files from Medicare linked with patient

Electronic Health Records (EHR) information. This dataset was created to study medically complex diabetic patients, thus patients were included if they satisfied the following two conditions: (C1) they were selected by a validated algorithm for identifying diabetic patients via claims data, and (C2) they were medically homed with an established plurality provider algorithm at participating provider group (a large Midwestern multi-specialty provider). Patient data was included for every 90-day quarter from 2003–2013 in which they were alive on the first day of the quarter, had continuous Medicare Part A & B fee-for-service, and met the medical home criteria (C1) and (C2) ($n = 9101$).

Longitudinal information is regularly collected on diabetes patients during primary care visits and is therefore available in the EHR (Dassow 2007). For example, in our study dataset, patients who have undergone a treatment change or failed to meet glycemic goals are recommended to have their HbA1c value measured quarterly (American Diabetes Association 2018). We seek to use this information to construct a risk prediction model for uncontrolled HbA1c. This is in contrast to predictive models of short-term (e.g., three months) control constructed from data collected over short time intervals (Lee et al. 2002; Weisner et al. 2003) or methods that seek to predict long-term control using a single landmark event (Parast, Cheng, and Cai 2012).

CONTACT Ying-Qi Zhao  yqzhao@fredhutch.org  Public Health Sciences Division, Fred Hutchinson Cancer Research Center, Seattle, WA 98109.

 Supplementary materials for this article are available online. Please go to www.tandfonline.com/r/JASA.

© 2021 American Statistical Association

Our goal is to construct a high-quality risk prediction model that incorporates accumulating longitudinal information without subjecting patients to undue burden. Using full-length records, that is, those observed over the longest allowable time horizon, for each individual may lead to better prediction, but repeated biomarker collection, such as measuring HbA1c values, can be costly and inconvenient. Furthermore, in applying such a predictive model, there is an inherent trade-off between waiting for enough information to accumulate so as to make a high-quality prediction and rapid identification of a high-risk patient so as to maximize the effectiveness of preventative care. We propose a new approach, termed reinforced risk prediction, which recursively incorporates longitudinal information into a personalized risk prediction model for uncontrolled HbA1c. This model accounts for costs associated with delaying patient classification to a subsequent follow-up time at which more information would be available to improve the accuracy of a prediction. At each follow-up time, the proposed model categorizes a patient trajectory as: (i) classified as high-risk for uncontrolled HbA1c; (ii) classified as low-risk for uncontrolled HbA1c; and (iii) unclassified due to insufficient information. Patients classified as high- or low-risk are then treated accordingly, whereas unclassified patients are monitored to be classified at a later time point when more information will be available. Because each additional measurement incurs new cost (measured as treatment burden, resource expenditure, risk of an adverse event, or some combination of these factors), we optimize classification accuracy subject to a constraint on the total cost.

The proposed approach is related to the classification-with-reject-option framework in the classification literature in which the classifier is allowed to forgo making a prediction for a small cost (Chow 1970; Herbei and Wegkamp 2006; Bartlett and Wegkamp 2008; Yuan and Wegkamp 2010). However, unlike our setting, in the classification-with-reject-option framework, predictions are not made over time with accruing information nor is it required that every subject eventually be classified. Trapeznikov and Saligrama (2013) considered applying reject options in a multi-stage setting. However, they required specification of a constant penalty for each reject option, which is not meaningful in our context.

In addition, EHR data pose significant challenges that prevent direct application of existing methods for sequential classification. First, EHR data are often incomplete as a patient's information is recorded only if and when they visit a clinic. In our data, although protocol stipulated that patients were to measure their HbA1c every quarter (or at least every 6 months if their HbA1c is controlled), 37.8% of patients measured over a period of 11 years had ≤ 6 measurements in total. Moreover, the HbA1c measurement times are irregularly spaced and can vary across patients. Thus, existing methods such as Trapeznikov and Saligrama (2013) which require the features to be fully observed across all stages cannot be applied. Second, there are missing values for the response variable. Patients might dropout or die before the end of the follow-up period. Selecting only individuals with non-missing outcome values may introduce selection bias. These issues are well-documented in the literature on EHR data (Hripcsak et al. 2011; Hersh et al. 2013; Flood et al. 2015; McVeigh et al. 2016). To accommodate irregularly spaced data,

we use functional generalized linear models (Ratcliffe, Heller, and Leader 2002; Müller and Stadtmüller 2005; Li, Wang, and Carroll 2010) which perform functional principal component analysis (FPCA) (Ramsay and Silverman 2005) on longitudinal HbA1c trajectories, and then uses the leading functional principal component (FPC) scores to predict the binary response variable. In addition, we propose a weighting scheme to account for missing values in the response variable.

The remainder of the article is organized as follows. In Sections 2 and 3, we present the notation and formally describe the proposed reinforced risk prediction estimator. In Section 4, we present theoretical results characterizing the operating characteristics of the proposed method. The results of extensive simulation studies are presented in Section 5. In Section 6, we illustrate the reinforced risk prediction model with the diabetes data to identify patients who are more likely to have uncontrolled HbA1c level in the long-term. Section 7 concludes with a discussion. Technical results are relegated to the supplemental material.

2. Method: Reinforced Risk Prediction

2.1. Notation and Overview

We assume that the observed data are $\{(Z_i, \tilde{X}_{ij}, U_{ij}, Y_i), j = 1, \dots, m, i = 1, \dots, n\}$, which comprise n independent and identically distributed trajectories, one per subject, each of the form $\{(Z, \tilde{X}_j, U_j, Y), j = 1, \dots, m\}$, where: $Z \in \mathbb{R}^p$ denotes baseline patient information, $\tilde{X}_j \in \mathbb{R}$ denotes a longitudinal measurement taken at time $U_j \in [0, T]$ for $j = 1, \dots, m$, and $Y \in \{0, 1\}$ is a terminal outcome. Thus, the number of visits, m , and the times at which they occur, are random variables. We further assume that the longitudinal measurements, $\tilde{X}_j$, are noisy surrogates of a latent process, $X(t), t \in [0, T]$, so that $\tilde{X}_j = X(U_j) + \epsilon(U_j)$, where $\epsilon(\cdot)$ is a mean-zero white noise process that is independent of $X(\cdot)$ and $U_j, j = 1, \dots, m$ that has constant variance $\text{var}\{\epsilon(t)\} \equiv \sigma_\epsilon^2$.

The goal is to predict Y using accruing longitudinal information while continuing to take measurements only if the improvement to predictive accuracy outweighs the cost of delaying making a prediction. To do this, we develop what we term a reinforced risk prediction procedure as follows. Let $0 = t_0 < t_1 < \dots < t_K = T$ denote a set of time points at which we are allowed to make predictions. These time points need not coincide with the times at which longitudinal measurements are made. Let $\mathcal{S}_k = \{Z, X(t), : t \leq t_k\}$ denote baseline and longitudinal information up to and including time t_k . At each time point $t_k, k = 1, \dots, K - 1$, we consider a decision rule with a reject option, $d_k : \mathcal{S}_k \mapsto \{0, 1, \text{'no decision'}\}$, so that $d_k(\mathcal{S}_k) \in \{0, 1\}$ indicates that a definite decision was made, and thus follow-up measurements of the biomarker are not needed beyond t_k . The “no decision” option (reject option) is selected when there is not enough information, and we want to delay the decision to a later time. At the final time point ($t_K = T$), a classification must be given, and a standard classifier, $d_K : \mathcal{S}_K \mapsto \{0, 1\}$, is used. Let k^* denote the time point at which a definite decision is made. That is, $k^* = \min\{k : d_k(\mathcal{S}_k) = 0 \text{ or } 1\}$, then the ultimate prediction is $d_{k^*}(\mathcal{S}_{k^*})$.

2.2. An Optimal Sequential Decision Rule

Our goal is to balance the tradeoff between minimizing prediction error, which generally decreases with the amount of data collected, against the cost, burden, and risk associated with delaying prediction. The misclassification error associated with a sequential decision rule, $d = (d_1, \dots, d_K)$, is $M(d) = E[I\{Y \neq d_{k^*}(S_{k^*})\}]$. Let c_k denote the incremental cost incurred during the period $(t_{k-1}, t_k]$, $k = 1, \dots, K$. We recommend to consult with clinicians to determine the appropriate values for c_k 's. The average cost incurred by the sequential decision rule, d , is thus $C(d) = E\left(\sum_{k=1}^{k^*} c_k\right)$, where the expectation is over the distribution of k^* induced by the decision rule d .

Given a budget, B , for the average total cost, we want to minimize $M(d)$ while satisfying the constraint $C(d) \leq B$. Using the method of Lagrange multipliers, we instead consider the unconstrained problem of finding the minimizer of the penalized objective $\mathcal{L}_\lambda(d) = M(d) + \lambda C(d)$. We discuss the equivalence between the constrained and the penalized problem in [Appendix A](#). To compute a solution, we apply a variant of regression-based approximate dynamic programming (Bather 2000). To provide intuition for this estimator, we first derive its population analog.

For each $k = 1, \dots, K$, define $\rho_k^*(S_k) = P(Y = 1|S_k)$. The Bayes classifier for a decision to be made at time k given patient information S_k is

$$b_k^*(S_k) = \begin{cases} 1 & \text{if } \rho_k^*(S_k) > 0.5 \\ 0 & \text{otherwise.} \end{cases}$$

Suppose that at the final time point K (at which point a classification must be given), a yet unclassified patient presents with history S_K . The optimal decision rule is clearly

$$d_K^*(S_K) = b_K^*(S_K).$$

Define $J_K^\lambda(S_K) = \lambda c_K + P\{Y \neq d_K^*(S_K)|S_K\}$ to be the cost incurred at time K under d_K^* given patient history S_K . For a yet unclassified patient presenting at time $K-1$ with S_{K-1} , the expected cost associated with a decision rule d_{K-1} is

$$\begin{aligned} J_{K-1}^{\lambda, d_{K-1}}(S_{K-1}) &= \lambda c_{K-1} + E[I\{Y \neq d_{K-1}(S_{K-1})\} I\{d_{K-1}(S_{K-1}) \\ &\neq \text{"no decision"}\} | S_{K-1}] \\ &\quad + E[J_K^\lambda(S_K) I\{d_{K-1}(S_{K-1}) = \text{"no decision"}\} | S_{K-1}] \\ &= \lambda c_{K-1} + P\{Y \neq d_{K-1}(S_{K-1}) | S_{K-1}\} I\{d_{K-1}(S_{K-1}) \\ &\neq \text{"no decision"}\} \\ &\quad + E[J_K^\lambda(S_K) | S_{K-1}] I\{d_{K-1}(S_{K-1}) = \text{"no decision"}\}, \end{aligned}$$

from which it can be seen that the rule d_{K-1}^* that minimizes $J_{K-1}^{\lambda, d_{K-1}}$ is given by

$$d_{K-1}^*(S_{K-1}) = \begin{cases} b_{K-1}^*(S_{K-1}) & \text{if } P\{Y \neq b_{K-1}^*(S_{K-1}) | S_{K-1}\} \\ & \leq E[J_K^\lambda(S_K) | S_{K-1}] \\ \text{"no decision"} & \text{otherwise,} \end{cases}$$

which is equivalent to

$$d_{K-1}^*(S_{K-1}) = \begin{cases} 1 & \text{if } P\{Y = 1 | S_{K-1}\} \\ & \geq 1 - E[J_K^\lambda(S_K) | S_{K-1}] \\ 0 & \text{if } P\{Y = 1 | S_{K-1}\} \\ & \leq E[J_K^\lambda(S_K) | S_{K-1}] \\ \text{"no decision"} & \text{otherwise.} \end{cases}$$

Define $J_{K-1}^\lambda(S_{K-1}) = J_{K-1}^{\lambda, d_{K-1}^*}(S_{K-1})$; recursively for $k = K-2, \dots, 1$ define

$$d_k^{\lambda*}(S_k) = \begin{cases} b_k^*(S_k) & \text{if } P\{Y \neq b_k^*(S_k) | S_k\} \\ & \leq E[J_{k+1}^\lambda(S_{k+1}) | S_k] \\ \text{"no decision"} & \text{otherwise,} \end{cases}$$

where,

$$\begin{aligned} J_k^\lambda(S_k) &= \lambda c_k + P\{Y \neq b_k^*(S_k) | S_k\} I\{d_k^{\lambda*}(S_k) \neq \text{"no decision"}\} \\ &\quad + E[J_{k+1}^\lambda(S_{k+1}) | S_k] I\{d_k^{\lambda*}(S_k) = \text{"no decision"}\} \end{aligned}$$

is the expected cost for a yet unclassified patient presenting at time k with S_k who is to be classified under the optimal sequential decision rule in future. The preceding arguments, along with the principle of dynamic programming, establish the following theorem; additional details are provided in [Appendix B](#).

Theorem 1. Let $\lambda \geq 0$ be fixed. Assume that all requisite expectations exist and let $d^{\lambda*} = (d_1^{\lambda*}, \dots, d_K^{\lambda*})$ be as defined previously. Then $M(d^{\lambda*}) + \lambda C(d^{\lambda*}) \leq M(d) + \lambda C(d)$ for any other sequential decision rule d .

3. Estimation of an Optimal Sequential Decision Rule

[Theorem 1](#) shows that the optimal sequential decision rule is determined by: the Bayes classifiers, b_k^* , for $k = 1, \dots, K$; the conditional means of the cost functions $\delta_k^\lambda(S_k) \triangleq E[J_{k+1}^\lambda(S_{k+1}) | S_k]$ for $k = 1, \dots, K-1$; and the tuning parameter $\lambda \geq 0$. We describe estimation of the Bayes classifiers and conditional means for a fixed value of λ and then describe a tuning procedure for λ to ensure the estimated decision rule, $\hat{d}^\lambda$, (approximately) satisfies the original cost constraint $C(\hat{d}^\lambda) \leq B$. In [Appendix C](#), we also discuss procedures to account for missing values in the response variable and to perform variable selection when the baseline covariates Z are of high-dimensional.

3.1. Estimation of an Optimal Sequential Decision Rule for Fixed λ

In this section, we provide the estimation procedure for $d_k^{\lambda*}$. This estimator comprises a sequence of estimators for the Bayes classifiers b_k^* , $k = 1, \dots, K$ and models for $\delta_k^\lambda(S_k)$, $k = 1, \dots, K-1$. Note that the misclassification error under the Bayes rule at time k is $P\{Y \neq b_k^*(S_k) | S_k\} = \rho_k^*(S_k) \wedge \{1 - \rho_k^*(S_k)\}$; we will use a plug-in estimator of this error rate in constructing the thresholds for optimal decision rules. To simplify the notation, we omit λ in the following two subsections.

3.1.1. Estimation of b_k^*

Often longitudinal biomarker data are irregularly spaced subject to missingness; thus, we estimate the optimal classification rule using methods from functional data analyses (FDA) (Ramsay and Silverman 2005; Yao, Müller, and Wang 2005). We use functional principal components analysis (FPCA) to extract

meaningful features from each individual's trajectory which we use in constructing an estimator of $P(Y = 1|\mathcal{S}_k)$, and hence $b_k^*(\mathcal{S}_k)$, for each $k = 1, \dots, K$.

For each subject $i = 1, \dots, n$ and stage $k = 1, \dots, K$, define $\mathcal{S}_{i,k} = \{X_i(t), t \in [0, t_k], \mathbf{Z}_i\}$. We assume the following functional generalized linear model (Ratcliffe, Heller, and Leader 2002; Müller and Stadtmüller 2005):

$$\text{logit}\{P(Y_i = 1|\mathcal{S}_{i,k})\} = \beta_0^{(k)} + \int_0^{t_k} \beta^{(k)}(t)\{X_i(t) - \mu(t)\}dt + \mathbf{Z}_i^\top \boldsymbol{\gamma}^{(k)}, \quad (1)$$

where $\text{logit}(u) = \log\{u/(1-u)\}$, $\mu(t) = E\{X(t)\}$, $\beta_0^{(k)}$, $\boldsymbol{\gamma}^{(k)}$ are unknown coefficients, and $\beta^{(k)}(\cdot)$ is an unknown coefficient function that describes the association between Y_i and $X_i(t)$, $t \in [0, t_k]$. For $0 \leq t, t' \leq t_k$, let $\Sigma_X^{(k)}(t, t') = \text{cov}\{X_i(t), X_i(t')\}$ be the covariance function of the longitudinal biomarker up to time t_k . By Mercer's theorem (Mercer 1909), $\Sigma_X^{(k)}(t, t')$ can be decomposed as $\sum_{l=1}^{\infty} \lambda_l^{(k)} \phi_l^{(k)}(t) \phi_l^{(k)}(t')$, where $\boldsymbol{\phi}^{(k)} = \{\phi_l^{(k)} : l = 1, \dots, \infty\}$ are orthonormal eigenfunctions satisfying $\int_0^{t_k} \{\phi_l^{(k)}(t)\}^2 dt = 1$, $\int_0^{t_k} \phi_l^{(k)}(t) \phi_{l'}^{(k)}(t) dt = 0$, for $l \neq l'$, and $\lambda_1^{(k)} \geq \lambda_2^{(k)} \geq \dots$ are the corresponding eigenvalues. Based on this decomposition, a Karhunen-Loève expansion (Karhunen 1947) for $X_i(t)$, $t \in [0, t_k]$ is

$$X_i(t) = \mu(t) + \sum_{l=1}^{\infty} \xi_{il}^{(k)} \phi_l^{(k)}(t),$$

where $\xi_{il}^{(k)} = \int_0^{t_k} \{X_i(t) - \mu(t)\} \phi_l^{(k)}(t) dt$ are uncorrelated random variables with mean 0 and variance $\lambda_l^{(k)}$. In the literature, $\boldsymbol{\phi}^{(k)} = \{\phi_l^{(k)} : l = 1, \dots, \infty\}$ are called functional principal components (FPCs), and $\boldsymbol{\xi}_i^{(k)} = \{\xi_{il}^{(k)}, l = 1, \dots, \infty\}$ are called FPC scores for the i th individual. Notice that there are K different models corresponding to K time points, indexed by the superscript k . For different time points, the Karhunen-Loève expansions are also different, yielding different FPCs and FPC scores.

Express $\beta^{(k)}(t)$ in the associated eigenbasis to obtain $\beta^{(k)}(t) = \sum_{l=1}^{\infty} \beta_l^{(k)} \phi_l^{(k)}(t)$. Plugging the expressions for $X_i(t)$ and $\beta^{(k)}(t)$ back into (1), and using the fact that $\boldsymbol{\phi}^{(k)}$ form an orthonormal basis, it follows that (1) is equal to

$$\text{logit}\{P(Y_i = 1|\mathcal{S}_{i,k})\} = \beta_0^{(k)} + \sum_{l=1}^{\infty} \beta_l^{(k)} \xi_{il}^{(k)} + \mathbf{Z}_i^\top \boldsymbol{\gamma}^{(k)}. \quad (2)$$

In practice, one might select the first several leading FPCs to approximate $X_i(t)$, that is, $X_i(t) \approx \mu(t) + \sum_{l=1}^L \xi_{il}^{(k)} \phi_l^{(k)}(t)$, where L is the minimum number of components needed to explain a specified percentage of the variation, for example, 99% is a common choice. Thus,

$$\begin{aligned} \text{logit}\{P(Y_i = 1|\mathcal{S}_{i,k})\} &\approx \beta_0^{(k)} + \sum_{l=1}^L \beta_l^{(k)} \xi_{il}^{(k)} + \mathbf{Z}_i^\top \boldsymbol{\gamma}^{(k)} \\ &= \beta_0^{(k)} + (\boldsymbol{\xi}_i^{(k)})^\top \boldsymbol{\beta}^{(k)} + \mathbf{Z}_i^\top \boldsymbol{\gamma}^{(k)}, \end{aligned}$$

in which $\boldsymbol{\xi}_i^{(k)} = (\xi_{i1}^{(k)}, \dots, \xi_{iL}^{(k)})^\top$ and $\boldsymbol{\beta}^{(k)} = (\beta_1^{(k)}, \dots, \beta_L^{(k)})^\top$. The parameters $(\beta_0^{(k)}, \boldsymbol{\beta}^{(k)}, \boldsymbol{\gamma}^{(k)})$ can be estimated using the following two steps:

- Step 1. The FPC scores $\boldsymbol{\xi}_i^{(k)}$ are not known because the latent process $X_i(\cdot)$ is not directly observable, the eigenfunctions $\{\phi_l^{(k)}(\cdot)\}_{l \geq 1}$ are not known, or both. For the i th individual, define $m_i^{(k)} = \sup\{j : U_{ij} \leq t_k\}$, to be the index of the time point U_{ij} immediately preceding t_k . When the repeated observations are sufficiently dense for each subject, that is, if all $m_i^{(k)}$ are larger than some order of n (Zhang and Wang 2016), one can pre-smooth the discrete observations $\{(U_{ij}, \tilde{X}_{ij}) : j = 1, \dots, m_i^{(k)}\}$ by fitting a local linear regression, and the smoothed trajectories $\{\hat{X}_i^{(k)}(t) : 0 \leq t \leq t_k, i = 1, \dots, n\}$ are used to construct the covariance, eigenvalues/basis and FPC scores; details are given in Appendix D. When the design points are moderate or sparse, we use the principal components analysis through conditional expectation (PACE) approach to estimate the FPC scores (Yao, Müller, and Wang 2005). Let $\tilde{\mathbf{X}}_i^{(k)} = (\tilde{X}_{i1}, \dots, \tilde{X}_{im_i^{(k)}})^\top$, $\boldsymbol{\phi}_{il}^{(k)} = (\phi_l^{(k)}(U_{i1}), \dots, \phi_l^{(k)}(U_{im_i^{(k)}}))^\top$, $\boldsymbol{\mu}_i^{(k)} = (\mu(U_{i1}), \dots, \mu(U_{im_i^{(k)}}))^\top$, we obtain the best linear predictors for $\xi_{il}^{(k)}$ given the noisy observations $\tilde{\mathbf{X}}_i^{(k)}$ as $E[\xi_{il}^{(k)} | \tilde{\mathbf{X}}_i^{(k)}] = \lambda_l^{(k)} (\boldsymbol{\phi}_{il}^{(k)})^\top \Sigma_{\tilde{\mathbf{X}}_i^{(k)}}^{-1} (\tilde{\mathbf{X}}_i^{(k)} - \boldsymbol{\mu}_i^{(k)})$, where $\Sigma_{\tilde{\mathbf{X}}_i^{(k)}} = \text{cov}(\tilde{\mathbf{X}}_i^{(k)}, \tilde{\mathbf{X}}_i^{(k)})$. Plugging in estimates of $\lambda_l^{(k)}$, $\boldsymbol{\phi}_{il}^{(k)}$, $\Sigma_{\tilde{\mathbf{X}}_i^{(k)}}$, $\boldsymbol{\mu}_i^{(k)}$ then yields the estimated FPC scores $\hat{\xi}_{il}^{(k)} = \hat{E}[\xi_{il}^{(k)} | \tilde{\mathbf{X}}_i^{(k)}]$, $\hat{\boldsymbol{\xi}}_i^{(k)} = (\hat{\xi}_{i1}^{(k)}, \dots, \hat{\xi}_{im_i^{(k)}}^{(k)})^\top$. (See Appendix E for details.) Note here that the FPC scores are estimated using the noisy surrogates $\{\tilde{X}_{ij}, i = 1, \dots, n, j = 1, \dots, m_i^{(k)}\}$, regardless of dense design or not.
- Step 2. Fit a logistic regression of Y_i given $(\hat{\boldsymbol{\xi}}_i^{(k)}, \mathbf{Z}_i)$ to obtain the parameter estimates $(\hat{\beta}_0^{(k)}, \hat{\boldsymbol{\beta}}^{(k)}, \hat{\boldsymbol{\gamma}}^{(k)})$, where $\hat{\boldsymbol{\beta}}^{(k)} = (\hat{\beta}_1^{(k)}, \dots, \hat{\beta}_L^{(k)})^\top$.

The estimator of $P(Y_i = 1|\mathcal{S}_{i,k})$ is thus

$$\hat{P}(Y_i = 1|\mathcal{S}_{i,k}) = \frac{\exp\{\hat{\beta}_0^{(k)} + (\hat{\boldsymbol{\xi}}_i^{(k)})^\top \hat{\boldsymbol{\beta}}^{(k)} + \mathbf{Z}_i^\top \hat{\boldsymbol{\gamma}}^{(k)}\}}{1 + \exp\{\hat{\beta}_0^{(k)} + (\hat{\boldsymbol{\xi}}_i^{(k)})^\top \hat{\boldsymbol{\beta}}^{(k)} + \mathbf{Z}_i^\top \hat{\boldsymbol{\gamma}}^{(k)}\}},$$

from which we obtain $\hat{b}_k^*(\mathcal{S}_k) = 1$ if $\hat{P}(Y = 1|\mathcal{S}_k) > 0.5$ and 0 otherwise.

3.1.2. Estimation of $E\{J_{k+1}(\mathcal{S}_{k+1})|\mathcal{S}_k\}$

The plug-in estimator of $J_K(\mathcal{S}_K)$ is $\hat{J}_K(\mathcal{S}_K) = \lambda c_K + \hat{\rho}_K(\mathcal{S}_K) \wedge \{1 - \hat{\rho}_K(\mathcal{S}_K)\}$. We estimate $\delta_{K-1}(\mathcal{S}_{K-1}) \triangleq E\{J_K(\mathcal{S}_K)|\mathcal{S}_{K-1}\}$ by regressing $\hat{J}_K(\mathcal{S}_K)$ on $\mathcal{S}_{K-1}$ using a functional linear model. Let $\hat{\delta}_{K-1}(\mathcal{S}_{K-1})$ denote the estimated regression (details are given shortly) and $\hat{d}_{K-1}(\mathcal{S}_{K-1})$ the plug-in estimator of $d_{K-1}^*(\mathcal{S}_{K-1})$, thus

$$\begin{aligned} \hat{J}_{K-1}(\mathcal{S}_{K-1}) &= \lambda c_{K-1} + \left[\hat{\rho}_{K-1}(\mathcal{S}_{K-1}) \wedge \{1 - \hat{\rho}_{K-1}(\mathcal{S}_{K-1})\} \right] \\ &\quad \times I\{\hat{d}_{K-1}(\mathcal{S}_{K-1}) \neq \text{"no decision"}\} \\ &\quad + \hat{\delta}_{K-1}(\mathcal{S}_{K-1}) I\{\hat{d}_{K-1}(\mathcal{S}_{K-1}) = \text{"no decision"}\}. \end{aligned}$$

We estimate $\delta_{K-2}(S_{K-2}) \triangleq E\{J_{K-1}(S_{K-1})|S_{K-2}\}$ by regressing $\hat{J}_{K-1}(S_{K-1})$ on S_{K-2} using a functional linear model. This process is repeated recursively for $k = K-2, \dots, 1$ to obtain estimators $\hat{\delta}_1(S_1), \dots, \hat{\delta}_{K-1}(S_{K-1})$ and subsequently the plug-in estimators $\hat{d}_1, \dots, \hat{d}_K$.

For $k = 1, \dots, K-1$, we posit a functional linear model of the form

$$\delta_k(S_k) = \alpha_0^{(k)} + \int_0^{t_k} \alpha^{(k)}(t)\{X(t) - \mu(t)\}dt + \mathbf{Z}^\top \boldsymbol{\omega}^{(k)}, \quad (3)$$

where we have ignored the constant term λc_k (which can be absorbed into an intercept), $\alpha_0^{(k)}, \boldsymbol{\omega}^{(k)}$ are the unknown parameters, and $\alpha^{(k)}(\cdot)$ is an unknown parameter function. Write $\alpha^{(k)}(t)$ as a linear combination of FPCs, $\alpha^{(k)}(t) = \sum_{l=1}^{\infty} \alpha_l^{(k)} \phi_l^{(k)}(t)$. After choosing the first L leading FPCs, the right-hand side of Equation (3) is approximated by

$$\alpha_0^{(k)} + (\boldsymbol{\xi}^{(k)})^\top \boldsymbol{\alpha}^{(k)} + \mathbf{Z}^\top \boldsymbol{\omega}^{(k)},$$

where $\boldsymbol{\xi}^{(k)}$ are the FPC scores and $\boldsymbol{\alpha}^{(k)} = (\alpha_1^{(k)}, \dots, \alpha_L^{(k)})^\top$. We estimate the unknown parameters $\alpha_0^{(k)}, \boldsymbol{\alpha}^{(k)}$, and $\boldsymbol{\omega}^{(k)}$ using least squares with outcome $\hat{J}_{k+1}(S_{k+1})$ and covariates $(\hat{\boldsymbol{\xi}}^{(k)}, \mathbf{Z})$ for $k = 1, \dots, K-1$. An estimator of $\delta_k(S_k)$ is thus $\hat{\delta}_k(S_k) = \hat{\alpha}_0^{(k)} + (\hat{\boldsymbol{\xi}}^{(k)})^\top \hat{\boldsymbol{\alpha}}^{(k)} + \mathbf{Z}^\top \hat{\boldsymbol{\omega}}^{(k)}$.

3.2. Cross-Validation to Choose Optimal λ

For a new individual indexed by $i = n+1$. At time t_k , $k = 1, \dots, K$, let $\mathcal{S}_{n+1,k} = (\hat{\boldsymbol{\xi}}_{n+1}^{(k)}, \mathbf{Z}_{n+1})$ denote the features for this new individual. The estimated FPC score $\hat{\boldsymbol{\xi}}_{n+1}^{(k)}$ can be computed based on the longitudinal data $\tilde{\mathbf{X}}_{n+1} = (\tilde{X}_{n+1,1}, \dots, \tilde{X}_{n+1,m_{n+1}^{(k)}})^\top$ and the FPCA objects obtained from the training data.

For a fixed λ , a decision rule for this new individual has the following form. At time points $k = 1, \dots, K-1$,

$$\hat{d}_k^\lambda(S_{n+1,k}) = \begin{cases} 1 & \text{if } \hat{P}(Y=1|\mathcal{S}_{n+1,k}) \geq 1 - \hat{\delta}_k^\lambda(S_{n+1,k}) \\ 0 & \text{if } \hat{P}(Y=1|\mathcal{S}_{n+1,k}) \leq \hat{\delta}_k^\lambda(S_{n+1,k}) \\ \text{"no decision"} & \text{otherwise.} \end{cases}$$

At the last time point $t_K = T$,

$$\hat{d}_K(S_{n+1,K}) = \begin{cases} 1 & \text{if } \hat{P}(Y=1|\mathcal{S}_{n+1,K}) > 0.5 \\ 0 & \text{otherwise,} \end{cases}$$

where the value of λ does not affect $\hat{d}_K$. When λ is large, the "no decision" option is rarely invoked, meaning that predictions are based on the information from time interval $[0, t_1]$ only. Alternatively, when λ is small, the "no decision" option is used frequently and thus predictions are made using longer patient trajectories.

For a prespecified budget limit B , we want to find the minimum λ such that its corresponding sequential decision rule $\hat{d}^\lambda = (\hat{d}_1^\lambda, \dots, \hat{d}_{K-1}^\lambda, \hat{d}_K)$ satisfy $C(\hat{d}^\lambda) \leq B$. Write $\hat{d}_k^\lambda = (\hat{d}_1^\lambda, \dots, \hat{d}_k^\lambda)$, $k = 1, \dots, K-1$, and let $H_k(\hat{d}_{k-1}^\lambda)$ be a binary

state variable indicating whether a definite decision was made before time t_k ($1 = \text{no}, 0 = \text{otherwise}$). It follows that

$$C(\hat{d}^\lambda) = E\left[c_1 + \sum_{k=1}^{K-1} c_{k+1} H_k(\hat{d}_{k-1}^\lambda) I\{\hat{d}_k^\lambda(S_k) = \text{"no decision"}\}\right].$$

We use cross-validation to estimate an optimal λ . Specifically, we partition the study cohort randomly into Q folds that are roughly equal in size. Let V_q be the index set of the q th fold, $q = 1, \dots, Q$. For a fixed λ , and for each $q \in \{1, \dots, Q\}$, we use the q th fold as a validation set and combine the other folds into a training set. From the training set, we fit FPCA decompositions and obtain the parameter estimates. Then we use the fitted model to make the prediction for individuals in the validation set. The average cost on the q th fold is

$$\hat{C}^{(q)}(\hat{d}^\lambda) = c_1 + \frac{1}{n_q} \sum_{i \in V_q} \left[\sum_{k=1}^{K-1} c_{k+1} H_k(\hat{d}_{k-1}^\lambda) I\{\hat{d}_k^\lambda(S_{i,k}) = \text{"no decision"}\} \right],$$

where n_q is the number of individuals in V_q . The cross-validated estimator of $C(\hat{d}^\lambda)$ is $\hat{C}(\hat{d}^\lambda) = \frac{1}{Q} \sum_{q=1}^Q \hat{C}^{(q)}(\hat{d}^\lambda)$. We will preselect a range of candidate values for λ , and search for the minimum λ where the budget constraint is satisfied, that is, $\hat{C}(\hat{d}^\lambda) \leq B$.

4. Theoretical Results

In this section, we establish theoretical results for the estimated optimal sequential decision rule using the proposed method. The theoretical developments are nontrivial. First, we extend the recent work of Kong et al. (2016) from partially functional linear models to generalized functional partially linear models. In addition, our proposed procedure involves modeling two components that are intertwined, where we estimate $\delta_{K-1}(S_{K-1})$ by regressing $\hat{J}_K(S_K)$ on S_{K-1} , and $\hat{J}_K(S_K)$ itself is obtained from the fitted logistic regression model (1). Finally, because the best decision rule at the current stage depends on the best decision rules at later stages, the theoretical results need to be worked out stage by stage conditioning on the estimated rule at future stages. We consider the dense sampling design. In addition, we assume that there is no missing value for Y , and that the number of baseline covariates p is fixed. Technical conditions and detailed proofs are relegated to Appendix F to H.

We assume that the data-generating model satisfies Equation (1), and recall $\beta^{(k)}(t) = \sum_{l=1}^{\infty} \beta_l^{(k)} \phi_l^{(k)}(t)$. Let $\beta_{0*}^{(k)}, \boldsymbol{\gamma}_*^{(k)}$ and $\beta_{l*}^{(k)}$ denote the true values of the parameters for $l \geq 1$. Following Kong et al. (2016), we write $\tilde{\beta}_l^{(k)} = \{\lambda_l^{(k)}\}^{1/2} \beta_l^{(k)}$, so that the FPC scores are on a common scale of variability. In addition, let L_n be the number of FPC components used to approximate $X(t)$, and denote $\tilde{\boldsymbol{\theta}}^{(k)} = (\beta_0^{(k)}, \tilde{\boldsymbol{\beta}}^{(k)\top}, \boldsymbol{\gamma}^{(k)\top})^\top$ where $\tilde{\boldsymbol{\beta}}^{(k)} = (\tilde{\beta}_1^{(k)}, \dots, \tilde{\beta}_{L_n}^{(k)})^\top$. The proposed estimator is $\tilde{\boldsymbol{\theta}}^{(k)} = (\hat{\beta}_0^{(k)}, \tilde{\boldsymbol{\beta}}^{(k)\top}, \hat{\boldsymbol{\gamma}}^{(k)\top})^\top$ where $\tilde{\boldsymbol{\beta}}^{(k)} = (\tilde{\beta}_1^{(k)}, \dots, \tilde{\beta}_{L_n}^{(k)})^\top$, $\tilde{\beta}_l^{(k)} = \{\lambda_l^{(k)}\}^{1/2} \hat{\beta}_l^{(k)}$, $l = 1, \dots, L_n$. Let $\tilde{\boldsymbol{\theta}}_*$ denote the true value of $\tilde{\boldsymbol{\theta}}^{(k)}$, that is, $\tilde{\boldsymbol{\theta}}_*^{(k)} = (\beta_{0*}^{(k)}, \tilde{\boldsymbol{\beta}}_*^{(k)\top}, \boldsymbol{\gamma}_*^{(k)\top})^\top$ where $\tilde{\boldsymbol{\beta}}_*^{(k)} =$

$(\tilde{\beta}_{1*}^{(k)}, \dots, \tilde{\beta}_{L_n*}^{(k)})^\top, \tilde{\beta}_{l*}^{(k)} = \{\lambda_l^{(k)}\}^{1/2} \beta_{l*}^{(k)}, l = 1, \dots, L_n$. The following theorem establishes the convergence rates of $\tilde{\theta}^{(k)}$ and $\hat{P}(Y = 1|S_k)$ to $\tilde{\theta}_*^{(k)}$ and $P(Y = 1|S_k)$ respectively.

Theorem 2. Under conditions (A1)–(A9) in the [supplementary material](#), for any $k = 1, \dots, K$, we have $\|\tilde{\theta}^{(k)} - \tilde{\theta}_*^{(k)}\| = O_p(L_n^{1/2} n^{-1/2})$. Subsequently, $\hat{P}(Y = 1|S_k) - P(Y = 1|S_k) = O_p(L_n n^{-1/2})$.

Recall $\alpha^{(k)}(t) = \sum_{l=1}^{\infty} \alpha_l^{(k)} \phi_l^{(k)}(t)$ and let $\alpha_{0*}^{(k)}, \omega_*^{(k)}$ and $\alpha_{l*}^{(k)}, l \geq 1$ be the true values of the parameters in model (3). Similarly, we write $\tilde{\alpha}_l^{(k)} = \{\lambda_l^{(k)}\}^{1/2} \alpha_l^{(k)}$. Denote $\tilde{\mathbf{t}}^{(k)} = (\alpha_0^{(k)}, \tilde{\alpha}^{(k)\top}, \omega^{(k)\top})^\top$ where $\tilde{\alpha}^{(k)} = (\tilde{\alpha}_1^{(k)}, \dots, \tilde{\alpha}_{L_n}^{(k)})^\top$. The proposed estimator is $\check{\mathbf{t}}^{(k)} = (\hat{\alpha}_0^{(k)}, \check{\alpha}^{(k)\top}, \hat{\omega}^{(k)\top})^\top$ where $\check{\alpha}^{(k)} = (\check{\alpha}_1^{(k)}, \dots, \check{\alpha}_{L_n}^{(k)})^\top$. Let $\tilde{\mathbf{t}}_*^{(k)}$ denote the true value of $\tilde{\mathbf{t}}^{(k)}$, i.e., $\tilde{\mathbf{t}}_*^{(k)} = (\alpha_{0*}^{(k)}, \tilde{\alpha}_*^{(k)\top}, \omega_*^{(k)\top})^\top$ where $\tilde{\alpha}_*^{(k)} = (\alpha_{1*}^{(k)}, \dots, \alpha_{L_n*}^{(k)})^\top, \tilde{\alpha}_{l*}^{(k)} = \{\lambda_l^{(k)}\}^{1/2} \alpha_{l*}^{(k)}, l = 1, \dots, L_n$. The following theorem establishes the convergence rates of $\check{\mathbf{t}}^{(k)}, \hat{\delta}_k^\lambda(S_k)$ and thus the consistency results for the estimated decision rules.

Theorem 3. Under conditions (A1)–(A12) in the [supplementary material](#), for any $k = 1, \dots, K - 1$, we have $\|\check{\mathbf{t}}^{(k)} - \tilde{\mathbf{t}}_*^{(k)}\| = O_p(L_n^{1/2} n^{-1/2})$. Subsequently, $\hat{\delta}_k^\lambda(S_k) - \delta_k^\lambda(S_k) = O_p(L_n n^{-1/2})$. Furthermore, our estimated sequential decision rule is consistent. That is, for fixed λ , $(\hat{d}_1^\lambda, \dots, \hat{d}_{K-1}^\lambda, \hat{d}_K)$ converges in probability to the optimal decision rule $(d_1^{\lambda*}, \dots, d_{K-1}^{\lambda*}, d_K^*)$.

Remark 1. The convergence rates of $\hat{\delta}_k^\lambda(S_k)$ are comparable to those in Theorem 3.1 of Laber and Staicu (2018). Under partially functional linear models with treatment-covariate interactions, Laber and Staicu (2018) showed that the functional Q-learning estimator (an estimator of the linear predictor function) converges at a rate of $O_p(L_n n^{-1/2} + L_n^{1/2} n^{-\Delta})$, where Δ is dictated by the slowest rate of convergence among the estimators of μ, Σ_X , and σ_ϵ . In the dense design we consider here, it is possible to obtain $\Delta = 1/2$ (Zhang and Wang 2016), and thus $O_p(L_n n^{-1/2} + L_n^{1/2} n^{-\Delta}) = O_p(L_n n^{-1/2})$, the same rate as we obtained for $\hat{\delta}_k^\lambda(S_k)$.

5. Simulation Studies

We conducted extensive simulation studies to evaluate the finite sample performance of our proposed method. Similar to Goldsmith, Greven, and Crainiceanu (2013), we generated the longitudinal biomarker using the following model:

$$\tilde{X}_i(t) = X_i(t) + \epsilon_i(t) = \mu(t) + \sum_{l=1}^4 \xi_{il} \phi_l(t) + \epsilon_i(t), \quad (4)$$

for t on the equally spaced grid $\{s/60, s = 0, \dots, 60\}$, $\mu(t) = 0$ and $\epsilon_i(t) \sim N(0, \sigma_\epsilon^2)$. The eigenfunctions were chosen as $\phi = \{\phi_1(t) = 1; \phi_2(t) = \sqrt{3}(2t - 1); \phi_3(t) = \sqrt{5}(6t^2 - 6t + 1); \phi_4(t) = \sqrt{7}(20t^3 - 30t^2 + 12t - 1)\}$. The score variances are $\lambda_l = 0.75^{l-1}, l = 1, \dots, 4$, and the FPC scores are normally

distributed with $\xi_{il} \sim N(0, \lambda_l)$. The measurement error variance is $\sigma_\epsilon^2 = 0.01$.

Notice that for each individual, there are 61 time points in total. To reflect commonly encountered missing data situations, we assume that the curve $\tilde{X}_i(t)$ is incompletely observed. We randomly observe 10 points from grid $\{0, \dots, 24/60\}$, assuming that the baseline measurement is always available. The remaining 10 observations are randomly chosen from $\{25/60, \dots, 1\}$ without replacement. We pre-smooth the discrete observations, and then use the smoothed trajectories to construct the covariance, eigenvalues/basis and FPC scores. A design plot is also created, which displays the assembled pairs (U_{ij}, U_{ik}) of all subjects with $n = 100$ (Figure S1 of the [supplementary material](#)). As this figure illustrates, the assembled pairs are dense in the domain plane.

The outcome Y is generated according to four different scenarios:

Scenario 1. Let $Z \sim \text{Bernoulli}(0.6)$ represents a baseline covariate, we assume that

$$\text{logit}\{P(Y_i = 1|X_i(t), t \in [0, 1], Z_i)\} = a_0 + a_1 X_i(1) + a_2 Z_i,$$

where $a_0 = 0, a_1 = -0.5, a_2 = 0.5$. That is, the outcome Y_i depends only on $X_i(1)$, the biomarker value at the final time point.

Scenario 2. We assume that

$$\begin{aligned} \text{logit}\{P(Y_i = 1|X_i(t), t \in [0, 1], Z_i)\} \\ = a_0 + a_1 b_i X_i(1) + a_2 (1 - b_i) X_i(0.5) + a_3 Z_i, \end{aligned}$$

where $a_0 = 0, a_1 = a_2 = -0.5, a_3 = 0.5$, and $b_i \sim \text{Bernoulli}(0.5)$. Notice that b_i indicates whether Y_i depends on $X_i(1)$ or $X_i(0.5)$. Hence, for half of the population, the binary outcome depends on $X_i(1)$ only, and for the other half, it depends on $X_i(0.5)$ only.

Scenario 3. We assume that

$$\begin{aligned} \text{logit}\{P(Y_i = 1|X_i(t), t \in [0, 1], Z_i)\} \\ = a_0 + \int_0^1 \beta(t) X_i(t) dt + a_1 Z_i, \end{aligned} \quad (5)$$

where $a_0 = -0.5, a_1 = 0.5$ and $\beta(t) = \sqrt{2} \sin(\pi t/2) + 3\sqrt{2} \sin(3\pi t/2)$ (similar to the simulation example in Shin (2009)). In this scenario, the outcome is associated with the overall trajectory of the longitudinal biomarker.

Scenario 4. This scenario mimicks our motivating example, in which the outcome is whether the average HbA1c value was lower or higher than 7% in the 6th year, and the longitudinal biomarker is the HbA1c values over the first 5 years. We generated $\tilde{X}_i(t)$ according to model (4), but now on interval $[0, 1.2]$ with grid $\{s/60, s = 0, \dots, 72\}$. Here, the interval $[0, 1]$ corresponds to the first 5 years, and $[1, 1.2]$ corresponds to the 6th year. As before, $\tilde{X}_i(t)$ is observed with 10 points from $\{0, \dots, 24/60\}$ and another 10 points from $\{25/60, \dots, 72/60\}$. The outcome Y_i is 1 if the average of $\tilde{X}_i(t)$ during time interval $(1, 1.2]$ is negative, and coded as 0 otherwise. A binary indicator variable O_i is generated,

$$\begin{aligned} \text{logit}\{P(O_i = 1|X_i(t), t \in [0, 1], Z_i)\} \\ = -0.5 - \int_0^1 X_i(t) dt + 0.5 Z_i, \end{aligned}$$

where $O_i = 1$ indicates that the i th individual does not drop out or die during the interval $[0, 1]$ (the first 5 years), and $O_i = 0$ otherwise. The outcome Y_i is missing either because $O_i = 0$ or the i th individual has no measurements available in the interval $(1, 1.2]$. Here, Y_i is missing at random, and the missing rate is around 50%. We predict the binary outcome using the longitudinal biomarker information $\tilde{X}_i(t), t \in [0, 1]$ and other baseline covariates. We varied the measurement error variance σ_ϵ^2 among $(0.01, 0.1, 0.5)$ to investigate how it affected the simulation results.

The performance of the following four methods were compared:

1. Use the baseline biomarker value $\tilde{X}_i(0)$ to predict Y_i (Baseline). For this analysis, the cost was 1 as the biomarker is measured only at time 0.
2. The traditional functional generalized linear model (FGLM). For this analysis, the cost is 61, because each individual is followed for 61 time points. The estimated FPC scores based on these complete trajectories are used to fit the FGLM and make the predictions.
3. Wait until a fixed time point and used the up-to-date longitudinal information to make the prediction (Fixed-30, 35, 40, 45, 50, and 55). For instance, Fixed-30 indicates that we use data from the first 30 time points $\{0, 1/60, \dots, 29/60\}$ to fit FPCA, and then the estimated FPC scores are used for prediction.
4. The proposed reinforced risk prediction (Reinforced-30, 35, 40, 45, 50, and 55). For example, Reinforced-30 means that the average follow up time was no longer than 30 points, although the individual follow up times could vary. In other words, on average, we would like to make a definite prediction before time $t = 0.5$. In our simulation studies, the predictions were made at $\{t_1 = 24/60, t_2 = 25/60, \dots, t_{37} = 1\}$.

The proposed and competing methods were implemented using R software. The proposed method is freely available through the reinforcedPred package hosted on the comprehensive R network (cran.org). For each scenario, we considered three sample sizes for training datasets: $n = 100, 400$ or 1000 . The performances of four methods were evaluated on an independent test dataset of 10,000 individuals. The reported misclassification rate and average cost are estimated using 1000 Monte Carlo replications.

Results for Scenarios 1 to 3 are presented in Figure 1. As expected, the performance of different methods improved as sample sizes increased. The baseline analysis incurred little cost, but suffered from a large misclassification error. On the other hand, traditional functional generalized linear models led to a small misclassification error, but at a higher cost. The proposed method achieved a balance between these two methods, which produced a high-quality prediction within the cost constraint. In all three scenarios, Reinforced-45 had an error rate comparable to FGLM, but with significantly lower cost. Compared to Fixed-30, 35, $\dots$, 55, reinforced risk prediction consistently resulted in an equal or smaller misclassification error while yielding a lower average cost.

Figure 2 shows the results for Scenario 4. In this scenario, the outcome Y_i could be missing. Inverse probability weighting (IPW) was used, where the weights were fitted on $\tilde{X}_i(t), t \in [0, 1]$ and Z_i . When $\sigma_\epsilon^2 = 0.01$, Reinforced-45, 50, 55 and FGLM had similar misclassification errors. However, as σ_ϵ^2 increased to 0.5, there were notable differences among these methods. That is, the more noisier the biomarker data is, the preferable to wait longer before making a prediction. In practice, we could obtain an estimator of σ_ϵ^2 based on the training data (see step (2) in Appendix E for details), and then use the estimated value as a guideline.

We also evaluated the performance of our method with moderate-dimensional baseline covariates. We adopted the same data generating mechanisms as those in Scenarios 1, 2, and 3, but the dimension of baseline covariates was increased to 50. In particular, $\mathbf{Z} = (Z_1, \dots, Z_{50})$, where $Z_1, \dots, Z_{25}$ were independent and identically distributed with Bernoulli(0.6), and $Z_{26}, \dots, Z_{50}$ were independent and identically distributed with $N(0, 1)$. Among these baseline covariates, only Z_1 attributed to the outcome variable. The results are presented in Figure 3, where we employed the variable selection techniques in Appendix C. Again, the proposed method required less resources, but still made a high-quality prediction.

In addition, we report the parameter estimation results for the proposed method. Here we focus on Scenario 3 since its data-generating model (5) aligns with Equation (1) at $t_K = T$, that is, the outcome model is correctly specified. For model (5), Table 1 displays the bias, standard deviation, and mean squared error (MSE) of $\hat{\alpha}_0, \hat{\alpha}_1$; and the functional mean squared error (MSE_f) of $\hat{\beta}(t)$, $\text{MSE}_f = E \left[\int_0^1 \{\hat{\beta}(t) - \beta(t)\}^2 dt \right]$. The results show that the bias, MSE and MSE_f are reasonably small, and decrease as sample size increases.

Additional simulations are presented in Appendix I. In particular, we showed that even if the outcome Y was generated from a probit model and we used the logistic regression model (1) instead, the proposed method still performed well by providing a cost-effective prediction.

6. Applications to the Diabetes Study

We analyze a cohort of diabetes patients treated under UW Health, one of the country's largest physician group practices (DuGoff et al. 2018). This dataset contains 9101 patients who were enrolled from 2003 to 2013, and who were followed up every 3 months until the 4th quarter of 2013. Thus, the longest follow up time is 11 years (44 quarters) while the average follow up time is 4.57 years (18.28 quarters). We restricted our analysis to a subpopulation of 8635 patients who had at least one HbA1c measurement in the first 5 years. The binary outcome Y was defined as: "1" = average HbA1c is lower than or equal to 7% (under tight control) in the 6th year, "0" = otherwise. The longitudinal biomarkers were the HbA1c values in the first 5 years. The baseline covariates $\mathbf{Z}$ included gender, race, medication, baseline age, and 45 indicators for other comorbidities, such as indicators for congestive heart failure and chronic kidney disease. Race was categorized into 3 groups: White, Black, and

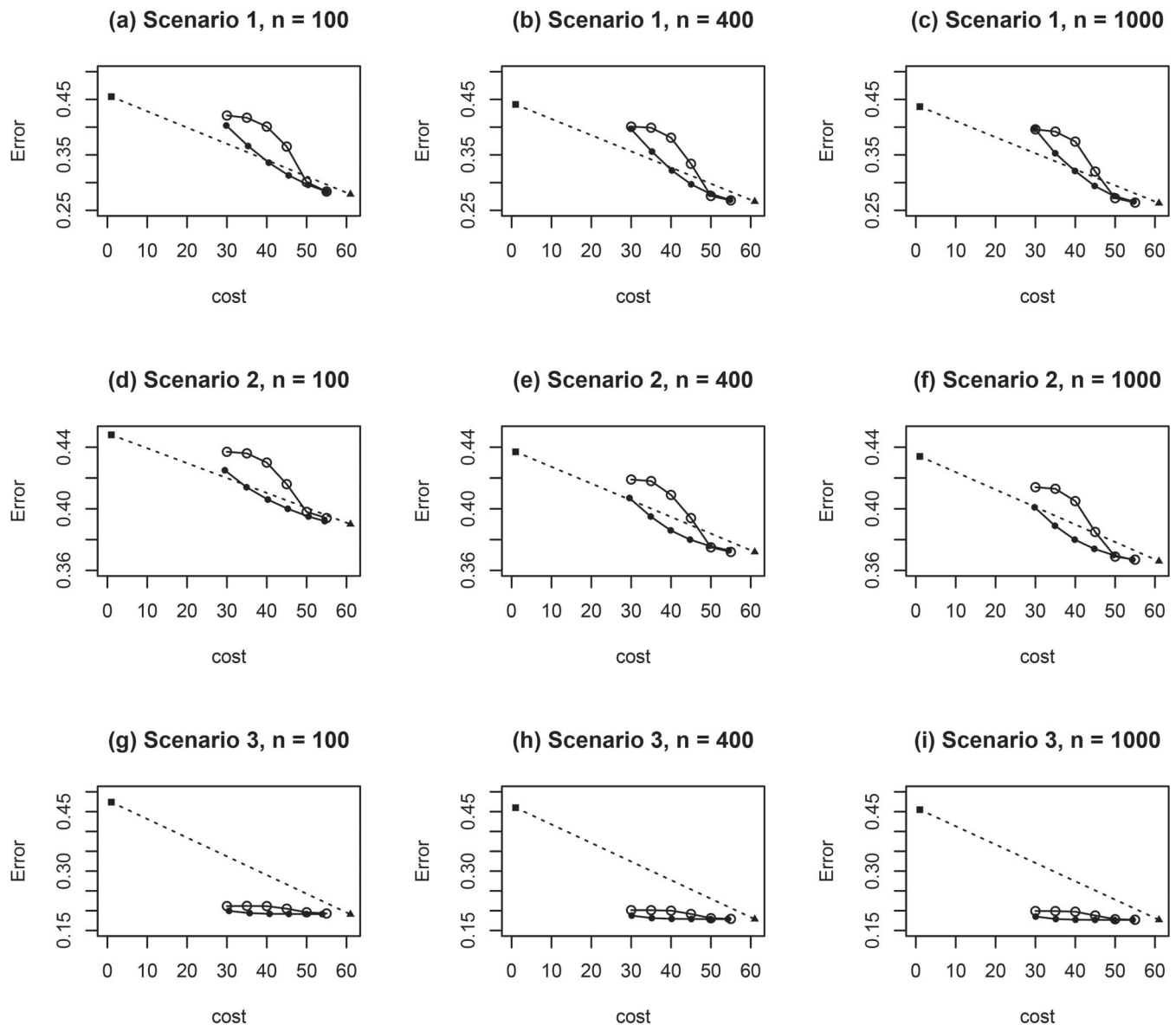


Figure 1. Misclassification error vs. cost for various risk prediction methods under Scenarios 1 to 3. The symbols are ■: Baseline, ▲: FGLM, ○: Fixed-30,35,...,55, ●: Reinforced-30,35,...,55.

other. Medication was coded as “1” = receiving either insulin, sulfonylurea or oral hypoglycemic agents, “0” = none of the above three.

Note that patients are mostly over the age of 65 at baseline, and are complex patients with comorbidities and complications (at later stages of diabetes). Hence, this cohort can be approximately considered as a homogeneous population with common mean and covariance functions of HbA1c levels over time. In addition, we create a design plot where assembled pairs (U_{ij}, U_{ik}) are displayed (Figure S2 of the supplementary material). While the data available per subject are sparse, the assembled data fill the domain of the covariance surface quite densely, and thus it is appropriate to use PACE.

As indicated in the previous sections, the outcome Y was missing for some participants as they either dropped out or died during the first 5 years, or simply did not have any HbA1c measurements available in the 6th year. In total, 3394 participants had at least one measurement in the 6th year. IPW techniques

from Appendix C were employed to address this problem. We performed FPCA on longitudinal HbA1c trajectories, and then fit a logistic regression model for the missing data mechanism, with leading FPC scores and baseline characteristics Z included as covariates.

We applied the comparison methods outlined in Section 5. There was a total of 21 time points counting the baseline quarter across 5 years. Hence, the cost was 21 in the FGLM method. Fixed-12, 13, ..., 20 and Reinforced-12, 13, ..., 20 were conducted. Here, Fixed-13 meant that we used the data from the first 13 time points, that is, data from the first 3 years including the baseline, to fit FPCA. Reinforced-13 indicated that on average, the decisions for patients were made by the end of the 3rd year, although the decision time for each patient might vary.

To evaluate the performance of different methods, we randomly split the data into a training and a testing set with 1:1 ratio. The risk prediction model was fitted using the training set, and then the prediction rule was applied to the testing set.

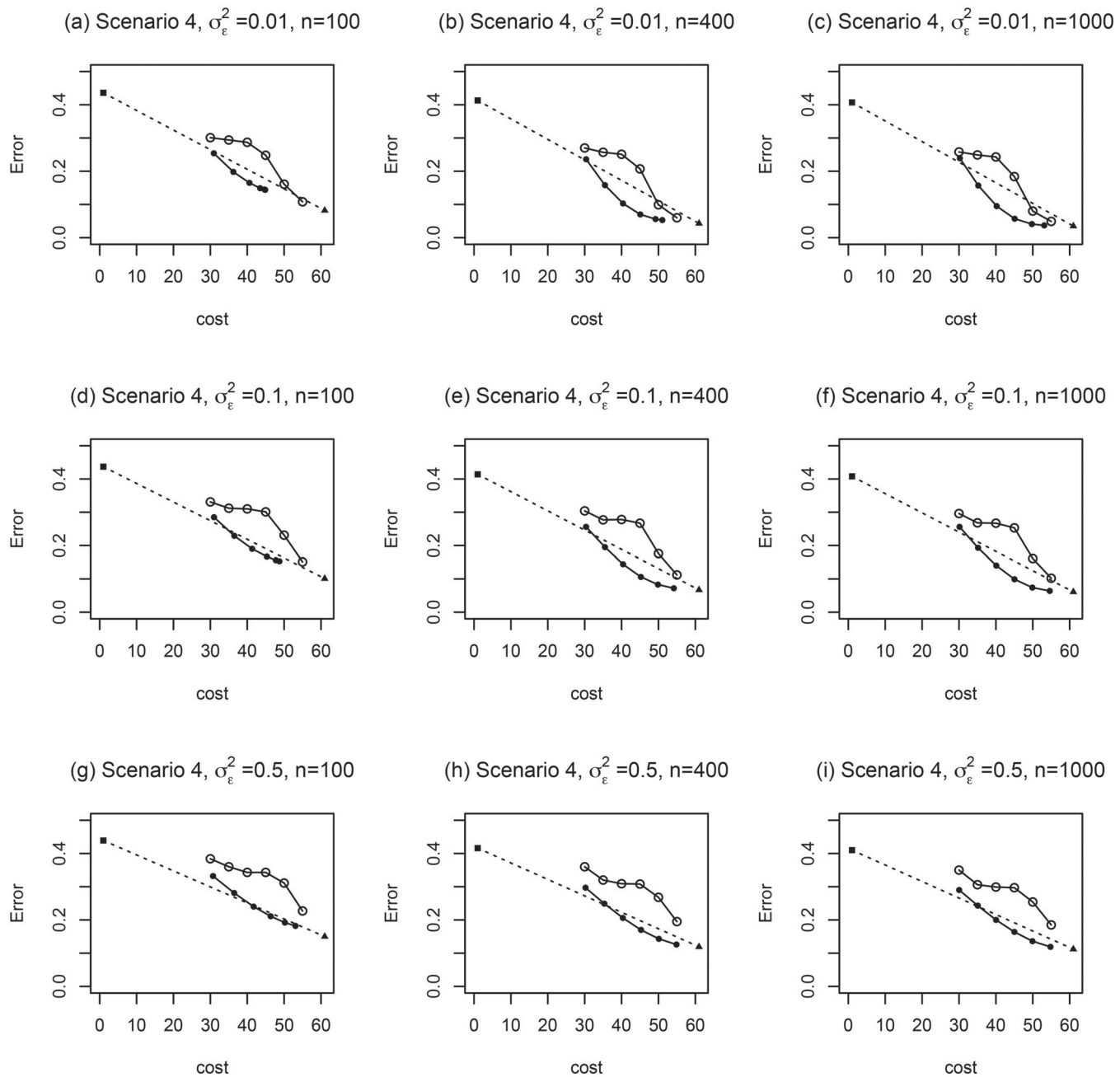


Figure 2. Misclassification error vs cost for various risk prediction methods under Scenario 4. The symbols are ■: Baseline, ▲: FGLM, ○: Fixed-30, 35, ..., 55, ●: Reinforced-30, 35, ..., 55.

The empirical misclassification rates and the average cost were calculated based on the testing set. This process was repeated multiple times, and we reported the mean of misclassification error and the average cost from 1000 replications in Figure 4. Note that Reinforced-18, 19, 20 correspond to the same $\lambda = 0$, and hence the dots are overlapped. The resulting error rate was 0.257 with the traditional functional generalized linear models (FGLM), which used the full data. The Reinforced-17 model, which on average made predictions one year earlier, yielded an error rate of 0.263. The Reinforced-13 model also produced a fairly close error rate of 0.285, while making predictions two years earlier on average compared with FGLM method. Hence, with the proposed method, we are able to make an early prediction, with little sacrifice in prediction accuracy. In addition,

the curve produced by the proposed method lies below the line connecting the results of the baseline and the FGLM, as well as the line produced by Fixed-12, 13, ..., 20. While the average costs of the Fixed-12, 13, ..., 20 are similar as those of the Reinforced-12, 13, ..., 20 methods, the proposed procedure can wisely allocate the time for prediction based on individual information and lead to lower error rates across all time points. This also indicates that our method is cost-effective.

7. Discussion

Motivated by risk prediction among complex diabetic patients we proposed reinforced risk prediction procedure as a means of predicting long-term risk while balancing the cost, which

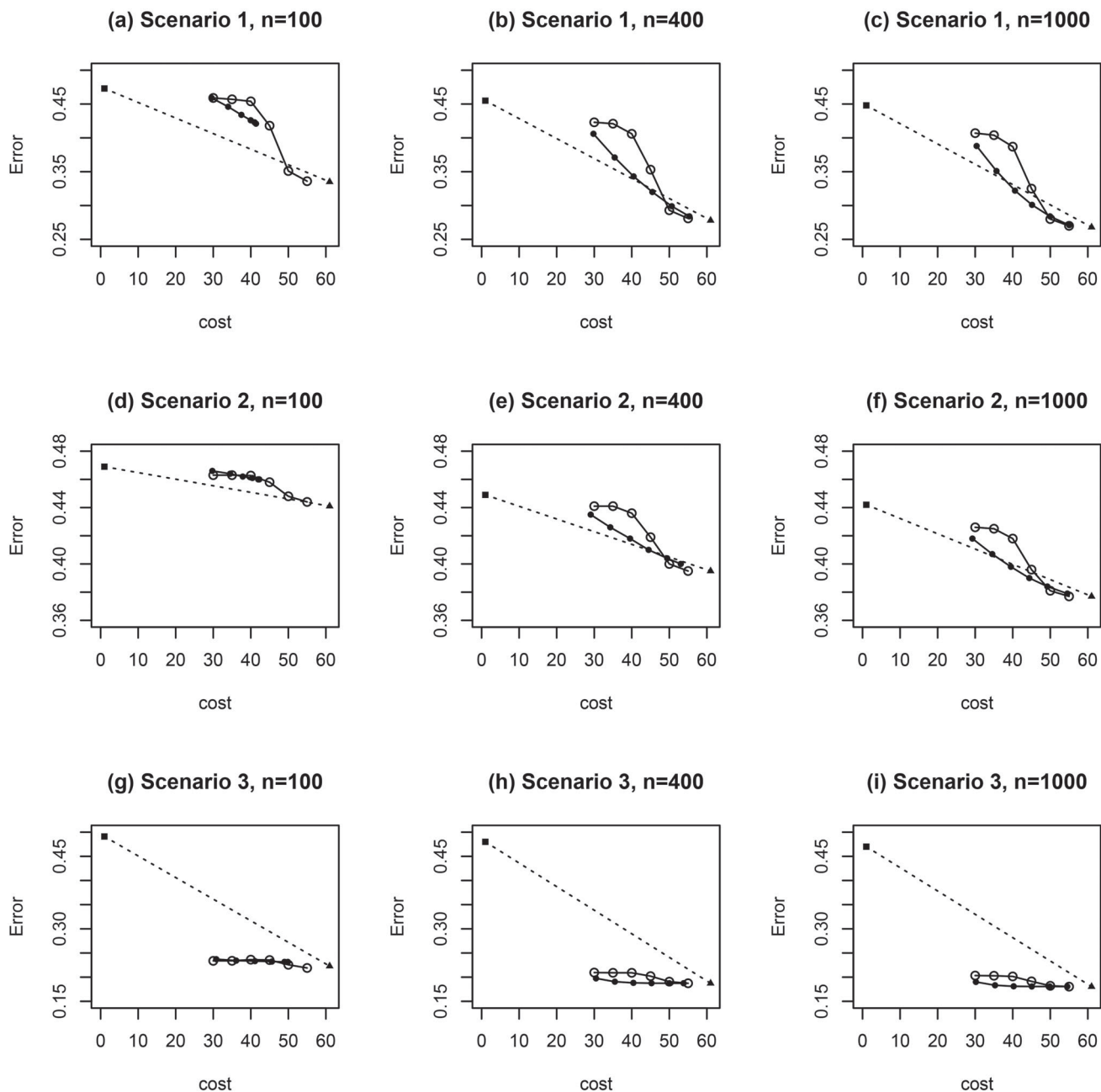


Figure 3. Misclassification error vs cost for various risk prediction methods under Scenarios 1 to 3 with 50 baseline covariates. The symbols are ■: Baseline, ▲: FGLM, ○: Fixed-30, 35, ..., 55, ●: Reinforced-30, 35, ..., 55.

can be interpreted as monetary cost or the cost of lost time to prevent complications. The proposed method, which relies on approximate dynamic programming and functional data analyses, is quite general and could be used to construct adaptive triage plans for a wide range of chronic diseases. In addition, it can be applied to other areas such as genomics, electronic commerce, and growth curve analysis, such as classifying temporal gene expression curves into known gene groups, predicting end prices of online auctions, and predicting the risk of being overweight.

Compared to FGLM which uses the full-length records, the proposed method saves cost but also sacrifices some prediction accuracy. Indeed, there is always benefit-cost tradeoffs in making clinical decisions, where cost can be monetary cost, patient

Table 1. Parameter estimation results under Scenario 3.

	$\hat{a}_0$			$\hat{a}_1$			$\hat{\beta}(t)$
	bias	SD	MSE	bias	SD	MSE	MSE_f
$n = 100$	-0.047	0.64	0.41	0.056	0.72	0.52	2.13
$n = 400$	-0.010	0.28	0.08	0.018	0.30	0.09	0.41
$n = 1000$	-0.009	0.18	0.03	0.006	0.19	0.03	0.22

burden, and side effects, etc. The interpretation and tolerance over the tradeoff will depend on patient/clinician preferences. Hence, when making the decision in practice once we obtain data-driven rules, we need to communicate with the patient and the clinician, and make plans accordingly.

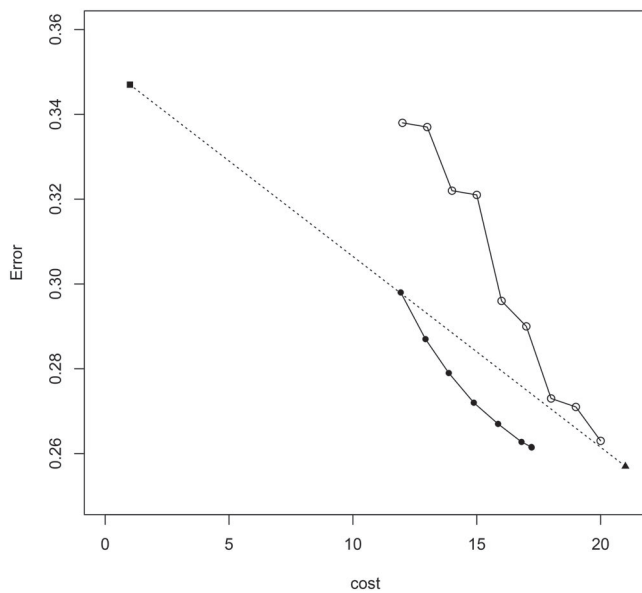


Figure 4. Analysis of the diabetes study data. Misclassification error vs cost for various risk prediction methods. The symbols are ■: Baseline, ▲: FGLM, ○: Fixed-12, ..., 20, ●: Reinforced-12, ..., 20.

In our current framework, we perform separate FPCA at different prediction time points. Suggested by one reviewer, it would be more efficient to borrow information across different prediction time points and implement an on-line update of the algorithm. This will require incremental updates in covariance function once new time points come in. One possibility is to impose certain parametric structure on the correlation function, for example, by assuming that $\text{corr}\{X_i(t), X_i(t')\}$ is a parametric function of $|t - t'|$. New algorithms and theoretical developments are needed, which are beyond the scope of this paper.

In the theoretical analysis, we assume that there is no missing value for Y , and that p is fixed. The cases with missing values of Y or the number of baseline covariates diverging, that is, $p_n \rightarrow \infty$ pose additional theoretical challenges. In particular, it would be interesting to investigate how inverse probability weighting affects the convergence rates of the parameter estimates. There are several other directions this work may be extended to. First, we can incorporate individual patient preferences (Butler et al. 2018) into the framework. Second, we can consider settings with multiple phases of treatment. In principle, reinforced risk prediction could be applied within each phase though alignment and delayed effects would have to be carefully considered. In addition, there could be multiple biomarkers or multi-category involved in the decision making. We are currently exploring solutions for these cases.

Supplementary Material

The supplementary material includes additional simulation results, regularity conditions, auxiliary results and proofs.

Funding

The authors gratefully acknowledge support by R01DK108073 awarded by the National Institutes of Health. Yinghao Pan's work was supported, in part, by funds provided by the University of North Carolina at Charlotte.

ORCID

Yinghao Pan <http://orcid.org/0000-0002-4022-1815>
Maureen A. Smith <http://orcid.org/0000-0003-4370-000X>

References

- ADVANCE Collaborative Group (2008), "Intensive Blood Glucose Control and Vascular Outcomes in Patients With Type 2 Diabetes," *New England Journal of Medicine*, 358, 2560–2572. [1090]
- American Diabetes Association (2018), "6. Glycemic Targets: Standards of Medical Care in Diabetes—2018," *Diabetes Care*, 41, S55–S64. [1090]
- Bartlett, P. L., and Wegkamp, M. H. (2008), "Classification With a Reject Option Using a Hinge Loss," *Journal of Machine Learning Research*, 9, 1823–1840. [1091]
- Bather, J. (2000), *Decision Theory: An Introduction to Dynamic Programming and Sequential Decisions*, Vol. 180, Chichester: Wiley. [1092]
- Butler, E. L., Laber, E. B., Davis, S. M., and Kosorok, M. R. (2018), "Incorporating Patient Preferences Into Estimation of Optimal Individualized Treatment Rules," *Biometrics*, 74, 18–26. [1100]
- Centers for Disease Control and Prevention (2017), *National Diabetes Statistics Report, 2017*, Atlanta, GA: Centers for Disease Control and Prevention, US Dept of Health and Human Services. [1090]
- Chien, K.-L., Lin, H.-J., Lee, B.-C., Hsu, H.-C., and Chen, M.-F. (2010), "Prediction Model for High Glycated Hemoglobin Concentration Among Ethnic Chinese in Taiwan," *Cardiovascular Diabetology*, 9, 59. doi:10.1186/1475-2840-9-59. [1090]
- Chow, C. (1970), "On Optimum Recognition Error and Reject Tradeoff," *IEEE Transactions on Information Theory*, 16, 41–46. [1091]
- Dassow, P. L. (2007), "Measuring Performance in Primary Care: What Patient Outcome Indicators do Physicians Value?" *The Journal of the American Board of Family Medicine*, 20, 1–8. [1090]
- DuGoff, E. H., Walden, E., Ronk, K., Palta, M., and Smith, M. (2018), "Can Claims Data Algorithms Identify the Physician of Record?" *Medical Care*, 56, e16–e20. [1096]
- Flood, T. L., Zhao, Y.-Q., Tomayko, E. J., Tandias, A., Carrel, A. L., and Hanrahan, L. P. (2015), "Electronic Health Records and Community Health Surveillance of Childhood Obesity," *American Journal of Preventive Medicine*, 48, 234–240. [1091]
- Goldsmith, J., Greven, S., and Crainiceanu, C. (2013), "Corrected Confidence Bands for Functional Data Using Principal Components," *Biometrics*, 69, 41–51. [1095]
- Herbei, R., and Wegkamp, M. H. (2006), "Classification With Reject Option," *Canadian Journal of Statistics*, 34, 709–721. [1091]
- Hersh, W. R., Weiner, M. G., Embi, P. J., Logan, J. R., Payne, P. R., Bernstam, E. V., Lehmann, H. P., Hripcsak, G., Hartzog, T. H., Cimino, J. J., and Saltz, J. H. (2013), "Caveats for the Use of Operational Electronic Health Record Data in Comparative Effectiveness Research," *Medical Care*, 51, S30–S37. [1091]
- Hripcsak, G., Knirsch, C., Zhou, L., Wilcox, A., and Melton, G. B. (2011), "Bias Associated With Mining Electronic Health Records," *Journal of Biomedical Discovery and Collaboration*, 6, 48–52. [1091]
- Karhunen, K. (1947). *Über lineare Methoden in der Wahrscheinlichkeitssrechnung*, (Vol 37), Annales Academiae Scientiarum Fennicae: Ser. A 1. Mathematica - Physica. [1093]
- Kong, D., Xue, K., Yao, F., and Zhang, H. H. (2016), "Partially Functional Linear Regression in High Dimensions," *Biometrika*, 103, 147–159. [1094]
- Kwon, H. N., Lee, Y. J., Kang, J.-H., Choi, J.-h., An, Y. J., Kang, S., Lee, D. H., Suh, Y. J., Heo, Y., and Park, S. (2014), "Prediction of Glycated Hemoglobin Levels at 3 Months After Metabolic Surgery Based on the 7-Day Plasma Metabolic Profile," *PLoS One*, 9, e109609. [1090]
- Laber, E. B., and Staicu, A.-M. (2018), "Functional Feature Construction for Individualized Treatment Regimes," *Journal of the American Statistical Association*, 113, 1219–1227. [1095]
- Lee, S. J., Klein, J. P., Barrett, A. J., Ringden, O., Antin, J. H., Cahn, J.-Y., Carabasi, M. H., Gale, R. P., Giral, S., Hale, G. A., et al. (2002). "Severity of Chronic Graft-Versus-Host Disease: Association with Treatment-Related Mortality and Relapse," *Blood*, 100, 406–414. [1090]

- Li, Y., Wang, N., and Carroll, R. J. (2010), "Generalized Functional Linear Models with Semiparametric Single-Index Interactions," *Journal of the American Statistical Association*, 105, 621–633. [1091]
- McVeigh, K. H., Newton-Dame, R., Chan, P. Y., Thorpe, L. E., Schreiberstein, L., Tatem, K. S., Chernov, C., Lurie-Moroni, E., and Perlman, S. E. (2016). "Can Electronic Health Records be Used for Population Health Surveillance? Validating Population Health Metrics Against Established Survey Data," *eGEMs (Generating Evidence & Methods to Improve Patient Outcomes)*, 4, 27. [1091]
- Mercer, J. (1909), "Functions of Positive and Negative Type, and their Connection the Theory of Integral Equations," *Philosophical Transactions of the Royal Society A*, 209, 415–446. [1093]
- Müller, H.-G., and Stadtmüller, U. (2005), "Generalized Functional Linear Models," *The Annals of Statistics*, 33, 774–805. [1091,1093]
- Parast, L., Cheng, S.-C., and Cai, T. (2012), "Landmark Prediction of Long-Term Survival Incorporating Short-Term Event Time Information," *Journal of the American Statistical Association*, 107, 1492–1501. [1090]
- Park, P., Griffin, S., Sargeant, L., and Wareham, N. (2002), "The Performance of a Risk Score in Predicting Undiagnosed Hyperglycemia," *Diabetes Care*, 25, 984–988. [1090]
- Ramsay, J. O., and Silverman, B. W. (2005), *Functional Data Analysis*, 2nd ed., New York: Springer-Verlag. [1091,1092]
- Ratcliffe, S. J., Heller, G. Z., and Leader, L. R. (2002), "Functional Data Analysis With Application to Periodically Stimulated Foetal Heart Rate Data. II: Functional Logistic Regression," *Statistics in Medicine*, 21, 1115–1127. [1091,1093]
- Shin, H. (2009), "Partial Functional Linear Regression," *Journal of Statistical Planning and Inference*, 139, 3405–3418. [1095]
- Trapeznikov, K., and Saligrama, V. (2013), "Supervised Sequential Classification Under Budget Constraints," in *Artificial Intelligence and Statistics*, pp. 581–589. [1091]
- Weisner, C., Thomas Ray, G., Mertens, J. R., Satre, D. D., and Moore, C. (2003), "Short-Term Alcohol and Drug Treatment Outcomes Predict Long-Term Outcome," *Drug and Alcohol Dependence*, 71, 281–294. [1090]
- Wilson, P. W., D'Agostino, R. B., Levy, D., Belanger, A. M., Silbershatz, H., and Kannel, W. B. (1998), "Prediction of Coronary Heart Disease Using Risk Factor Categories," *Circulation*, 97, 1837–1847. [1090]
- Yao, F., Müller, H.-G., and Wang, J.-L. (2005), "Functional Data Analysis for Sparse Longitudinal Data," *Journal of the American Statistical Association*, 100, 577–590. [1092,1093]
- Yuan, M., and Wegkamp, M. (2010), "Classification Methods With Reject Option Based on Convex Risk Minimization," *Journal of Machine Learning Research*, 11, 111–130. [1091]
- Zhang, X., and Wang, J.-L. (2016), "From Sparse to Dense Functional Data and Beyond," *The Annals of Statistics*, 44, 2281–2321. [1093,1095]