

CONTEXT-DRIVEN NEAREST NEIGHBOR IMPUTATION USING LANGUAGE REPRESENTATION

Anonymous authors

Paper under double-blind review

ABSTRACT

Missing data poses significant challenges for machine learning and deep learning algorithms. In this paper, we aim to enhance post-imputation performance, measured by machine learning utility (MLu). We introduce a nearest-neighbor-based imputation method, DrIM, designed for heterogeneous tabular datasets. However, calculating similarity in the data space becomes challenging due to the varying presence of missing entries across different columns. To address this issue, we leverage the representation learning capabilities of language models. By transforming the tabular dataset into a text-format dataset and replacing the missing entries with mask (or unk) tokens, we extract representations that capture contextual information. This mapping to a continuous representation space enables the use of well-defined similarity measurements. Additionally, we incorporate a contrastive learning framework to refine the representations, ensuring that the representations of observations with similar information in the observed columns, regardless of the missingness patterns, are closely aligned. To validate our proposed model, we evaluate its performance in missing data imputation across 10 real-world tabular datasets, demonstrating its ability to produce a Complete dataset having high MLu.

1 INTRODUCTION

Handling missing data is challenging, as machine learning (ML) and deep learning (DL) methods typically require complete datasets (Tan et al., 2013; Kingma & Welling, 2014; Goodfellow et al., 2014; Chen & Guestrin, 2016; Ke et al., 2017; Vaswani et al., 2017; Ho et al., 2020; An & Jeon, 2023). Rubin says that the goal of imputation tasks is to provide complete statistical inference (Rubin & Schenker, 1986). However, many recent imputation methods have focused on providing complete datasets that can yield high post-imputation performance, including machine learning utility (MLu). Ivanov et al. (2019); Yoon et al. (2018); Mattei & Frellsen (2019); Ipsen et al. (2021); Zhao et al. (2023); Du et al. (2024a); Wen et al. (2024) emphasize the importance of imputing missing data in terms of MLu. In this paper, we also focus on imputing incomplete datasets to provide a complete dataset that enables effective training of many ML and DL algorithms while achieving high MLu.

We empirically observed that nearest-neighbor-based imputation methods, such as k -nearest neighbors imputation (kNNI) (Troyanskaya et al., 2001) and its variants (Jiang & Yang, 2015; Thomas & Rajabi, 2021; Du et al., 2024b), can achieve high MLu. Moreover, Anil Jadhav & Ramanathan (2019); Keerin & Boongoen (2022); Ding et al. (2024) have demonstrated that these methods often outperform other imputation techniques. These nearest-neighbor-based methods impute missing values by leveraging the similarities between observations and utilizing the values of their nearest neighbors, thereby preserving local data characteristics (Troyanskaya et al., 2001; Tarsitano & Falcone, 2011).

However, nearest-neighbor-based imputation methods face two key challenges: 1) They struggle to effectively represent or address missing values. While various distance or similarity metrics have been proposed for missing data, their performance depends on how they handle these missing values (Juhola & Laurikkala, 2007; Muñoz & Hernández-González, 2012; AbedAllah & Shimshoni, 2016; Santos et al., 2020). For instance, assigning a distance of 0 when both values are missing can be inflexible and adversely affect the imputation process (Santos et al., 2020). 2) Applying these methods to heterogeneous (i.e., mixed-type) tabular data presents challenges. Jerez et al. (2010);

García-Laencina et al. (2010) have demonstrated that continuous and categorical variables exhibit an imbalance when measuring distance, with continuous variables having a more significant influence on the distance metric, affecting the results of imputation.

Therefore, we introduce a nearest-neighbor-based imputation method termed **DrIM** (Contextual-Driven Missing **IM**puter) to address these challenges. Our proposed method includes taking advantage of the language model’s ability to generate contextualized representations, which are learned from a large collection of text datasets. We utilize textual encoding to convert tabular data into text-format data and apply language models to obtain representations (Yin et al., 2020; Mei et al., 2021; Borisov et al., 2023; Radford & Narasimhan, 2018; Devlin et al., 2019; Nazir et al., 2023).

Our representation-based nearest-neighbor imputation method provides solutions to address challenges 1) and 2), respectively: 1) Missing values are represented as [MASK] (or [UNK]) tokens, allowing language models to utilize their representation capabilities and leverage contextual hints from the other columns. 2) By mapping both continuous and categorical columns into a continuous space using their representation vectors, we enable the use of well-defined similarity measures, such as cosine similarity. Furthermore, we incorporate fine-tuning using a contrastive learning framework to further enhance the language model’s representational capacity. This allows it to build more meaningful neighbors, which in turn improves the post-imputation performance.

Our primary contributions can be summarized as follows:

1. By adopting language models and transforming each record into a representation vector, we propose a nearest-neighbor-based imputation method that is applicable to heterogeneous tabular datasets.
2. We achieve state-of-the-art imputation performance in terms of MLu without any model training.
3. We provide an intuitive explanation while empirically confirming that our fine-tuning using contrastive learning enhances the MLu of the imputed dataset.

We validate the effectiveness of our proposed method by evaluating its imputation performance across 10 real-world tabular datasets, accounting for four missing data mechanisms and five missingness rates. We employ four MLu performance metrics, including classification tasks, model selection, and feature selection performance. Additionally, we demonstrate that applying a contrastive learning framework enhances imputation performance in terms of MLu.

2 RELATED WORKS

Imputation methods. The effectiveness of missing imputation ultimately lies in reducing bias and enhancing estimation efficiency in inference after imputation (Little & Rubin, 2019). This always depends on the missing data mechanism, and in general, a single imputation has the limitation of not accounting for the uncertainty induced by the imputation (van Buuren, 2012). Nevertheless, single imputation is widely used because of its convenience and its ability to improve estimation efficiency in downstream tasks (Álvarez Verdejo et al., 2021). In particular, a single imputation based on estimating conditional expectation performs well when the bias introduced by the imputation does not significantly affect the performance of the specific downstream task (Troyanskaya et al., 2001). Deep learning has garnered significant attention for its ability to flexibly model conditional expectations of observed data and effectively learn metrics in the observed data space (Yoon et al., 2018; Ivanov et al., 2019; Kyono et al., 2021; Wen et al., 2024). These two tasks are closely related to the nonparametric kNNI. kNNI is a method for calculating local conditional expectations and involves dimension reduction and metric learning to compute conditional expectations effectively.

Distance measures for missing data. kNNI is a widely used imputation method because it effectively utilizes the similarity between patterns to generate accurate estimates (Wilson & Martinez, 1997) and maintains the overall data distribution (Santos et al., 2017). However, its performance heavily relies on the chosen distance function. Several distance functions have been developed to handle heterogeneous tabular data and account for missing values. The HEOM (Aha et al., 1991; Wilson & Martinez, 1997) calculates distances based on feature types, using normalized Euclidean distance for continuous features and an overlap metric for categorical ones. Missing values are assigned a distance of 1, and if both are missing, 0. Juhola & Laurikkala (2007) refined HEOM to

maintain this approach. The HVDM (Stanfill & Waltz, 1986; Wilson & Martinez, 1997) functions similarly to HEOM but requires class target information for categorical features. SIMDIST (Muñoz & Hernández-González, 2012) introduces a heterogeneous similarity function incorporating prior knowledge, with similarity based on the presence or absence of values. Additionally, redefined versions of HVDM treat missing values as special cases (Santos et al., 2020), and the Mean Euclidean Distance (AbedAllah & Shimshoni, 2016) has been used for managing incomplete data in clustering. In contrast, our approach utilizes metric learning for imputing missing values. Instead of relying on existing distance functions, our method learns a distance (similarity) metric on a manifold using language models, specifically by computing the Euclidean distance between the representation vectors obtained from language models.

Notations. Suppose that an observation $\mathbf{x} \in \mathcal{X}_1 \times \cdots \times \mathcal{X}_p$ consists of both continuous and categorical variables (columns), where $\mathcal{X}_j$ denotes the support of j th variable. I_C and I_D represent the index sets for continuous and categorical variables, respectively, where $I_C \cup I_D = \{1, \dots, p\}$. $\mathbf{x}_j$ is a value of the j th column having its column name with C_j . Here, subscript j refers to the j th element. $g(\cdot; \theta)$ is a function that extracts the representation vector from the text input, where θ is its trainable parameter. The missingness pattern of an observation is defined by a corresponding missingness indicator vector $\mathbf{m} \in \{0, 1\}^p$, where $\mathbf{m}_j = 0$ if $\mathbf{x}_j$ is missing. An incomplete dataset is given $\{(\mathbf{x}^{(i)}, \mathbf{m}^{(i)})\}_{i=1}^n$ comprising n i.i.d. realizations of $\mathbf{x}$ and $\mathbf{m}$.

3 METHODOLOGY

DrIM¹ is a single imputation method based on the kNNI employing representations for each tubular record using language models. The main difference between the DrIM and the conventional kNNI approach lies in the way the distance metric between data points is learned using an underlying model for imputing missing values. This distance metric is defined as the Euclidean distance between the representation vectors of [CLS] tokens output by the BERT model (Devlin et al., 2019)². We will elucidate its theoretical properties in this section.

3.1 REPRESENTATION OF BERT

In this section, we will present some theoretical results demonstrating how the L_2 -distance between BERT representation vectors can reflect the similarity between observations. In the pre-training process of BERT, suppose that the pre-training dataset consists of paired sequences, $(\mathbf{t}, \mathbf{u}) \in \mathcal{T} \times \mathcal{U}$, where all paired sequences in $\mathcal{T} \times \mathcal{U}$ have the same fixed total length. Let $p(\mathbf{t}, \mathbf{u})$, $p(\mathbf{t})$, and $p(\mathbf{u})$ are joint and marginal distributions of $\mathbf{t}$ and $\mathbf{u}$, respectively. Let $\mathbf{e} \in \mathcal{U}$ be a sequence of the [PAD] tokens. We formally define the representation of BERT in the following definition.

Definition 1 (Representation of BERT). *Let $g(\cdot, \cdot; \theta) : \mathcal{T} \times \mathcal{U} \mapsto \mathbb{R}^D$ is a map of output vector located at the position of the [CLS] token in BERT. Then,*

1. *in the pre-training of BERT, the representation of BERT for $\mathbf{t} \in \mathcal{T}$ with respect to $\mathbf{u} \in \mathcal{U}$ is defined as $g(\mathbf{t}, \mathbf{u}; \theta)$;*
2. *in the process of DrIM for $\mathbf{t} \in \mathcal{T}$, the representation of BERT is defined as $g(\mathbf{t}, \mathbf{e}; \theta)$.*

In Definition 1, $\mathbf{u} \in \mathcal{U}$ represents auxiliary data used in the next sentence prediction task of BERT, which will be described later. This definition, while restrictive in that it fixes the dimensional sizes of $\mathbf{t}$ and $\mathbf{u}$ in BERT’s training, provides a convenient framework for analysis. In Definition 1, note that the sequence $\mathbf{u}$ is replaced with the sequence of [PAD] tokens in the process of DrIM.

Definition 2 (Pre-training of Next Sentence Prediction (NSP) via logistic regression). *In the pre-training of NSP, suppose that the positive and negative samples are sampled from the following distributions:*

$$\begin{aligned} (\mathbf{t}, \mathbf{u}) &\sim p(\mathbf{t}, \mathbf{u}) \quad (\text{positive sample}) \\ (\mathbf{t}, \mathbf{u}') &\sim p(\mathbf{t})p(\mathbf{u}') \quad (\text{negative sample}). \end{aligned}$$

¹The overall structure of DrIM is outlined in Figure 1.

²In this paper, we chose BERT for a comprehensive comparison and performance analysis; however, we have also included the experimental results using GPT-based models in Appendix A.8.

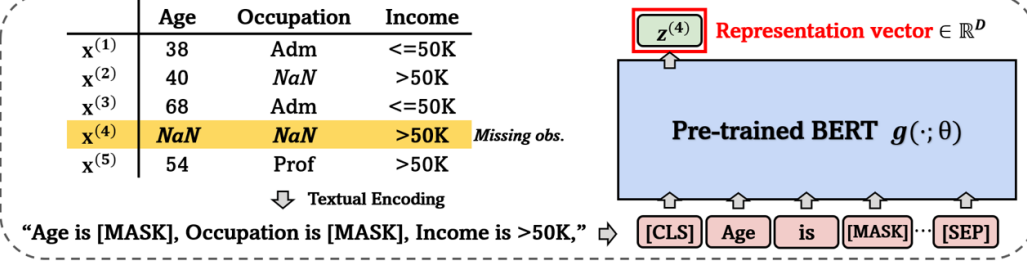
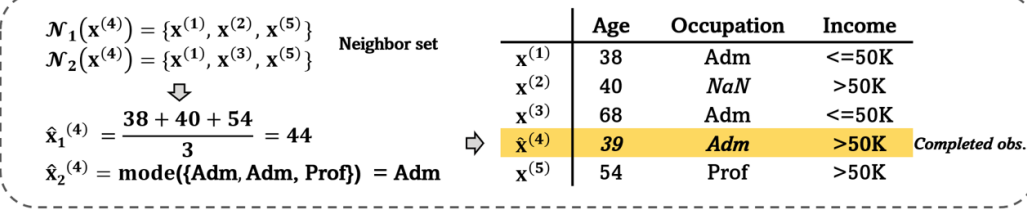
Phase1: Contextual Representation**Phase2: Context-Driven Imputation**

Figure 1: Overall structure of DrIM. In this case, we impute the missing values in $x^{(4)}$ using 3 nearest neighbors. The neighbor set varies since the presence of missing data differs by column.

Let a risk function

$$J_{LR}(\theta, \eta) = -\mathbb{E}_{p(\mathbf{t}, \mathbf{u})} \left[\log \frac{\exp(r(\mathbf{t}, \mathbf{u}; \theta, \eta))}{1 + \exp(r(\mathbf{t}, \mathbf{u}; \theta, \eta))} \right] - \mathbb{E}_{p(\mathbf{t})p(\mathbf{u}')} \left[\log \frac{1}{1 + \exp(r(\mathbf{t}, \mathbf{u}'; \theta, \eta))} \right],$$

where $r(\mathbf{t}, \mathbf{u}; \theta, \eta) = \Phi(g(\mathbf{t}, \mathbf{u}; \theta); \eta)$ and $\Phi(\cdot; \eta) : \mathbb{R}^D \mapsto \mathbb{R}$ is a scoring function parameterized with η . The training BERT for NSP is defined by estimating (θ, η) that minimizes an empirical version of J_{LR} .

In practice, BERT training is achieved by minimizing two loss functions: one for the masked language model (MLM) and the other for NSP. Here, we focus only on the loss function for NSP training, which can be analyzed within the theoretical framework of contrastive learning in unsupervised learning.

Theorem 1. Suppose that $\Phi(\cdot; \eta) : \mathbb{R}^D \mapsto \mathbb{R}$, the scoring function in Definition 2, is modeled as a linear function, i.e., $\forall \mathbf{z} \in \mathbb{R}^D, \Phi(\mathbf{z}; \eta) = \eta^\top \mathbf{z}$ and BERT is pre-trained by the NSP of Definition 2. Then, for two [CLS] tokens of $\mathbf{t}, \mathbf{t}' \in \mathcal{T}$ we have

$$\|g(\mathbf{t}, \mathbf{e}; \theta) - g(\mathbf{t}', \mathbf{e}; \theta)\| \geq \frac{1}{\|\eta\|} \cdot \left| \log \frac{p(\mathbf{e} | \mathbf{t})}{p(\mathbf{e} | \mathbf{t}')} \right|,$$

where p is the distribution defined on sequences from $\mathcal{T} \times \mathcal{U}$.

Theorem 1 demonstrates that the representation of BERT has the property that the L_2 -distance between representation vectors can reflect the similarity between observations. Using Bayes' rule, the log-likelihood ratio term in Theorem 1 can be re-written as $\log p(\mathbf{e} | \mathbf{t})/p(\mathbf{e} | \mathbf{t}') = \log \left(\frac{p(\mathbf{t}, \mathbf{e})/p(\mathbf{t}', \mathbf{e})}{p(\mathbf{t})/p(\mathbf{t}')} \right)$. As an input of BERT, the observed data is considered as pairs of $(\mathbf{t}, \mathbf{e}) \in \mathcal{T} \times \mathcal{U}$. Thus, the likelihood is written in terms of $p(\mathbf{t}, \mathbf{e})$, not $p(\mathbf{t})$. Therefore, the Euclidean distance between the representation vectors is an upper bound on the log-likelihood ratio with respect to the observed dataset, $\log p(\mathbf{t}, \mathbf{e})/p(\mathbf{t}', \mathbf{e})$, when compared to the marginalized log-likelihood ratio, $\log p(\mathbf{t})/p(\mathbf{t}')$. Theorem 1 provides a partially theoretical justification for defining neighbors based on the Euclidean distance between [CLS] tokens in a pre-trained BERT model.

3.2 PHASE1: CONTEXTUAL REPRESENTATION

To bridge an inherent gap between the structured nature of tabular data and the unstructured form of text data, we employ a textual encoding approach similar to Yin et al. (2020); Mei et al. (2021);

Borisov et al. (2023). By leveraging the encoding process of Definition 3, we obtain the textual dataset denoted as $\{\mathbf{t}^{(i)}\}_{i=1}^n$.

Definition 3 (Textual Encoding). *Each observation $\mathbf{x}^{(i)}$ is transformed into a sentence-like format $\mathbf{t}^{(i)}$. This transformation incorporates both the column names and their corresponding values, applying the following subject-predicate-object transformation:*

$$\begin{aligned}\mathbf{t}_j^{(i)} &= \begin{cases} "C_j \text{ is } \mathbf{x}_j^{(i)}, ", & \text{if } \mathbf{m}_j^{(i)} = 1 \\ "C_j \text{ is } [MASK], ", & \text{if } \mathbf{m}_j^{(i)} = 0 \end{cases}, \\ \mathbf{t}^{(i)} &= "[CLS]" \& \mathbf{t}_1^{(i)} \& \mathbf{t}_2^{(i)} \& \dots \& \mathbf{t}_p^{(i)} \& "[SEP]",\end{aligned}$$

where $\&$, called the ampersand operator, is used to concatenate one or more text strings to form a single text, and $[CLS]$ and $[SEP]$ token denote the start and end of the sequence, respectively.

Then, as in Definition 1, we obtain the representation vector corresponding to the observation $\mathbf{x}^{(i)}$, denoted as $\mathbf{z}^{(i)}$, as follows:

$$\mathbf{z}^{(i)} = g(\mathbf{t}^{(i)}, \mathbf{e}; \theta) \in \mathbb{R}^D,$$

where D denotes the dimension of the representation vector. The transformation process of contextual representation from $\mathbf{x}^{(i)}$ to $\mathbf{z}^{(i)}$ offers several advantages as follows:

1. Handling missing values with contextual information. Since we employ BERT, which is a MLM, we can represent missing values with $[MASK]$ tokens. This approach enables BERT to represent missing entries by leveraging contextual hints from other columns. It allows for the generation of a complete contextual representation of each observation without the need to disregard missing entries.

2. Missing-invariant mapping to continuous space. The language model $g(\cdot, \cdot; \theta)$ consistently outputs a D -dimensional real-valued vector for each observation, regardless of the various missingness patterns. By mapping to continuous representation space, we can utilize well-defined distances, such as the L_2 norm, between different observations.

3. Addressing heterogeneous tabular datasets. In the tabular domain, a central challenge lies in heterogeneity, as tabular data includes continuous, categorical, and even date (time) columns. Huang et al. (2020); Somepalli et al. (2021) address this challenge by processing continuous and categorical variables differently to obtain their representations. In contrast, we employ textual encoding to convert each observation into a sentence-like format, deriving representations from this sentence. This provides a type-agnostic preprocessing of the heterogeneous tabular dataset.

3.3 PHASE2: CONTEXT-DRIVEN IMPUTATION

Our proposed imputation procedure is based on the nearest neighbor set within the continuous representation space. The top- k nearest neighborhood on the representation space is defined as follows.

Definition 4. *The top- k nearest neighbors of $\mathbf{x}^{(i)}$, $i \in \{1, \dots, n\}$ with respect to the j th column are defined by the index set $\mathcal{N}_j(\mathbf{x}^{(i)}, k) \subset \{1, 2, \dots, n\}$, satisfying the following conditions:*

1. $|\mathcal{N}_j(\mathbf{x}^{(i)}, k)| = k$,
2. $\mathbf{m}_j^{(l)} = 1$ for all $l \in \mathcal{N}_j(\mathbf{x}^{(i)}, k)$, and
3. $\forall s \in \{1, \dots, n\} \setminus \mathcal{N}_j(\mathbf{x}^{(i)}, k)$ such that $\mathbf{m}_j^{(s)} = 1$

$$\cos(\mathbf{z}^{(i)}, \mathbf{z}^{(s)}) < \min_{l \in \mathcal{N}_j(\mathbf{x}^{(i)}, k)} \cos(\mathbf{z}^{(i)}, \mathbf{z}^{(l)}),$$

where $\cos(\cdot, \cdot)$ denotes the cosine similarity.

Note that the neighbor set includes only the indices of observations that are not missing in the dataset with respect to the j th column.

Our imputation procedure utilizes the neighbor set $\mathcal{N}_j(\mathbf{x}^{(i)}, k)$. For continuous columns, we impute using the average of the neighbor set. For categorical columns, we impute using the most frequent category among the neighbor set (see Algorithm 1 in Appendix for detailed procedure).

$$\hat{\mathbf{x}}_j^{(i)} = \begin{cases} \frac{1}{k} \sum_{l \in \mathcal{N}_j(\mathbf{x}^{(i)}, k)} \mathbf{x}_j^{(l)}, & \text{if } j \in I_C \\ \text{mode}(\mathbf{x}_j^{(l)} : l \in \mathcal{N}_j(\mathbf{x}^{(i)}, k)), & \text{if } j \in I_D \end{cases}.$$

Then, we obtain completed observation $\tilde{\mathbf{x}}$ as $\tilde{\mathbf{x}} = \mathbf{x} \odot \mathbf{m} + \hat{\mathbf{x}} \odot (\mathbf{1} - \mathbf{m})$, where $\odot$ denotes element-wise multiplication.

3.4 FINE-TUNING: CONTRASTIVE LEARNING

The population version of the objective function for fine-tuning is defined as follows (Ruderman et al., 2012; Belghazi et al., 2018):

$$\begin{aligned} & \max_{\theta} J_{CL}(\theta) \\ &= \mathbb{E}_{q(\mathbf{t}, \mathbf{u}, \mathbf{e})} [\cos(g(\mathbf{t}, \mathbf{e}; \theta), g(\mathbf{u}, \mathbf{e}; \theta))] - \log \mathbb{E}_{p(\mathbf{t}, \mathbf{e})p(\mathbf{u}, \mathbf{e})} [\exp(\cos(g(\mathbf{t}, \mathbf{e}; \theta), g(\mathbf{u}, \mathbf{e}; \theta)))], \end{aligned}$$

where $q(\mathbf{t}, \mathbf{u}, \mathbf{e})$ is the distribution of the positive sequence pair, and $p(\mathbf{t}, \mathbf{e})p(\mathbf{u}, \mathbf{e})$ is the distribution of the negative sequence pair. $q(\mathbf{t}, \mathbf{e}, \mathbf{u}) = q(\mathbf{u} | \mathbf{t}) \cdot p(\mathbf{t}, \mathbf{e})$ and the generating process of $\mathbf{u}$, i.e., $q(\mathbf{u} | \mathbf{t})$ is defined in Definition 5.

Definition 5 (Re-mask). Let ϕ be the re-masking function, which takes an anchor sample and a parameter r , specifying the number of columns to be masked. It randomly selects and masks r columns from the anchor sample, $\mathbf{t}$, to produce a positive sample, $\mathbf{u}$. Formally, the generating process of $\mathbf{u}$ is written as:

$$\begin{aligned} (\mathbf{t}, \mathbf{e}) &\sim p(\mathbf{t}, \mathbf{e}) \\ \mathbf{u} &= \phi(\mathbf{t}, r). \end{aligned}$$

The dropout augmentation strategy of Definition 5 is widely adopted to generate positive samples due to its simplicity and effectiveness (Chen et al., 2020; Miao et al., 2021; Wu et al., 2022; Jiang & Wang, 2022). For instance, if $r = 1$, we can obtain a positive sample of $\mathbf{t}^{(i)}$, which is denoted as $\mathbf{t}^{(i+)}$, by randomly selecting and masking one column from $\mathbf{t}^{(i)}$ as follows:

$$\begin{aligned} \mathbf{t}^{(i)} &= \text{"[CLS], } C_1 \text{ is } \mathbf{x}_1, C_2 \text{ is [MASK], } C_3 \text{ is } \mathbf{x}_3, C_4 \text{ is } \mathbf{x}_4, \text{[SEP]} \text{"} \\ \mathbf{t}^{(i+)} &= \text{"[CLS], } C_1 \text{ is } \mathbf{x}_1, C_2 \text{ is [MASK], } C_3 \text{ is } \mathbf{x}_3, C_4 \text{ is [MASK], [SEP]} \text{"}, \end{aligned}$$

where $\mathbf{t}^{(i+)}$ is sampled from $q(\mathbf{u} | \mathbf{t}^{(i)})$. In this case, the value of the C_4 column is re-masked.

Theorem 2. For θ such that the maximizer of $J_{CL}(\theta)$, $g(\mathbf{t}, \mathbf{e}; \theta)$ and $g(\mathbf{u}, \mathbf{e}; \theta)$ provide the representation vectors that maximize the lower bound of the mutual information between $\mathbf{t}$ and $\mathbf{u}$, where $(\mathbf{t}, \mathbf{u}) \sim q(\mathbf{u} | \mathbf{t}) \cdot p(\mathbf{t}, \mathbf{e})$.

Based on the InfoNCE loss calculation in the batch, in practice, we approximate $J_{CL}(\theta)$ as follows (Oord et al., 2018):

$$\begin{aligned} & \max_{\theta} \mathcal{L}(\theta) \\ &= \log \frac{\exp(\cos(g(\mathbf{t}^{(i)}, \mathbf{e}; \theta), g(\mathbf{t}^{(i+)}, \mathbf{e}; \theta)))}{\exp(\cos(g(\mathbf{t}^{(i)}, \mathbf{e}; \theta), g(\mathbf{t}^{(i+)}, \mathbf{e}; \theta))) + \sum_{l=1, l \neq i}^B \exp(\cos(g(\mathbf{t}^{(i)}, \mathbf{e}; \theta), g(\mathbf{t}^{(l)}, \mathbf{e}; \theta)))}, \end{aligned}$$

where B is the mini-batch size, $\mathbf{t}^{(i+)}$ is the positive sample of the i th sample ($i \in \{1, \dots, B\}$), and we define negative samples within a mini-batch (Tang et al., 2015; Mei et al., 2021). Note that, in the approximate of $J_{CL}(\theta)$ with $\mathcal{L}(\theta)$, $\mathbb{E}_{q(\mathbf{t}, \mathbf{u}, \mathbf{e})}$ is approximated with a single positive pair of $(\mathbf{t}^{(i)}, \mathbf{t}^{(i+)})$, and $\mathbb{E}_{p(\mathbf{t}, \mathbf{e})p(\mathbf{u}, \mathbf{e})}$ is approximated with the sequence pairs $(\mathbf{t}^{(i)}, \mathbf{t}^{(i+)})$, $(\mathbf{t}^{(i)}, \mathbf{t}^{(1)})$, $\dots$, $(\mathbf{t}^{(i)}, \mathbf{t}^{(B)})$.

Theorem 3. For θ such that the maximizer of $\mathcal{L}(\theta)$, $g(\mathbf{t}^{(i)}, \mathbf{e}; \theta)$ and $g(\mathbf{t}^{(i+)}, \mathbf{e}; \theta)$ provide the representation vectors that their cosine similarity is proportional to mutual information between $\mathbf{t}^{(i)}$ and $\mathbf{t}^{(i+)}$, where $\mathbf{t}^{(i+)} \sim q(\mathbf{u} | \mathbf{t}^{(i)})$.

We want to emphasize that Theorem 2 and 3 demonstrate that fine-tuning using contrastive learning allows the representations of BERT to be trained in such a way that they represent and maximize the mutual information between input sequences. In other words, our fine-tuning further enhances the representational capacity of the language model, capturing statistical dependencies between observed and missing entries in the *unseen tabular data scenario*. **Therefore, through fine-tuning with contrastive learning, we can perform missing data imputation using the fully observed data points that have high mutual information with the observations containing missing entries.**

4 EXPERIMENTS

4.1 OVERVIEW

We propose two versions of our model, DrIM: 1) DrIM_{BASE}, which utilizes the pre-trained BERT directly, and 2) DrIM_{FINE}, which incorporates contrastive learning framework³. While any BERT model can be used, we employ the BERT-base model⁴ to extract representation vectors.

Datasets. Similar to several recent studies (Mattei & Frellsen, 2019; Nazábal et al., 2020; Muzellec et al., 2020; Ipsen et al., 2021; Jarrett et al., 2022; Zhao et al., 2023), we utilize 10 real tabular datasets from the UCI repository⁵ and Kaggle⁶, which vary in sample sizes. The datasets are split into training and testing sets, with an 80% training and 20% testing ratio. Detailed statistics of these datasets are provided in Appendix A.4.

Additionally, following the approach in Muzellec et al. (2020); Jarrett et al. (2022); Zhao et al. (2023), we generate missing value masks for each dataset using three mechanisms (MCAR, MAR, MNAR) across four settings (MCAR, MAR, MNARL, MNARQ). Refer to Appendix A.7.1 for detailed descriptions and implementations of these mechanisms. We generate missing values at rates of 0.2, 0.3, 0.4, 0.6, and 0.8 for each mechanism.

Baseline models. In this experiment, we investigate our proposed method, DrIM, alongside 13 imputation methods, categorized into statistical imputation, ML-based imputation, and DL-based imputation as follows:

- Statistical imputation: Mean (Farhangfar et al., 2007), and kNNI (Troyanskaya et al., 2001)
- ML-based imputation: EM (Nelwamondo et al., 2007), SoftImpute (Mazumder et al., 2010), MICE (van Buuren & Groothuis-Oudshoorn, 2011), missForest (Stekhoven & Bühlmann, 2012), and Sinkhorn (Muzellec et al., 2020)
- DL-based imputation: GAIN (Yoon et al., 2018), VAEAC (Ivanov et al., 2019), MIWAE (Mattei & Frellsen, 2019), not-MIWAE (Ipsen et al., 2021), MIRACLE (Kyono et al., 2021), and ReMasker (Du et al., 2024a).

Detailed experimental settings for the reproducibility of these baseline models and our proposed model are provided in Appendix A.5.

Evaluation Metrics. To assess the imputation performance, we use two types of evaluation criteria: 1) imputation fidelity and 2) machine learning utility. The performance of imputers is reported by averaging the metrics across 10 real tabular datasets. Imputation fidelity is measured by the sum of the SMAPE (symmetric mean absolute percentage error) for continuous columns and the AR (accuracy error) for categorical columns, which we refer to as ARSMAPE, similar to (Miao et al., 2022). For machine learning utility, we adopt the metrics proposed by Hansen et al. (2023) to

³We conduct experiments using an NVIDIA A10 GPU, with our experimental codes implemented in PyTorch and scikit-learn.

⁴<https://huggingface.co/google-bert/bert-base-uncased>

⁵<https://archive.ics.uci.edu/>

⁶<https://www.kaggle.com/datasets/>

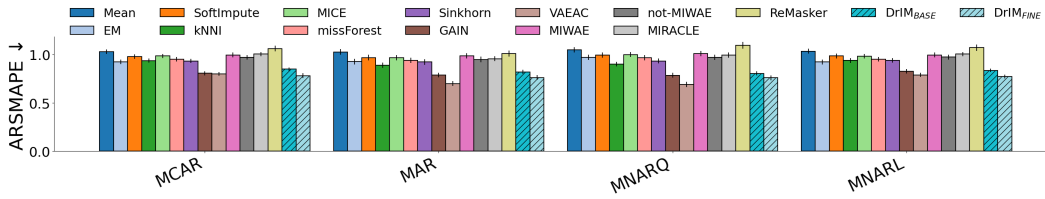


Figure 2: **Imputation fidelity** at 0.3 missingness rate. The missing mechanism corresponding to the figure is indicated below the figure. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported.

Table 1: **Machine learning utility** at 0.3 missingness under MAR and MNARL. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. The best value is bolded, and the second best is underlined.

model	MAR			MNARL		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.591 $\pm$.027	.548 $\pm$.041	.744 $\pm$.021	.560 $\pm$.027	.428 $\pm$.041	.797 $\pm$.017
kNNI	.601 $\pm$.027	.560 $\pm$.043	.868 $\pm$.014	.575 $\pm$.027	.419 $\pm$.049	.843 $\pm$.016
EM	.598 $\pm$.027	.531 $\pm$.043	.833 $\pm$.016	.580 $\pm$.027	.407 $\pm$.048	.804 $\pm$.019
SoftImpute	.590 $\pm$.027	.554 $\pm$.039	.845 $\pm$.014	.562 $\pm$.027	.342 $\pm$.045	.818 $\pm$.015
MICE	.598 $\pm$.027	.543 $\pm$.043	.859 $\pm$.014	.571 $\pm$.027	.339 $\pm$.050	.825 $\pm$.017
missForest	.596 $\pm$.027	.548 $\pm$.044	.833 $\pm$.015	.569 $\pm$.027	.366 $\pm$.048	.809 $\pm$.016
Sinkhorn	.599 $\pm$.027	.545 $\pm$.044	.845 $\pm$.015	.569 $\pm$.027	.377 $\pm$.049	.821 $\pm$.017
GAIN	.640 $\pm$.025	.602 $\pm$.029	.877 $\pm$.011	.621 $\pm$.025	.446 $\pm$.040	.843 $\pm$.016
VAEAC	.606 $\pm$.025	.612 $\pm$.018	.827 $\pm$.007	.598 $\pm$.025	.553 $\pm$.024	.816 $\pm$.012
MIWAE	.588 $\pm$.027	.519 $\pm$.048	.834 $\pm$.014	.556 $\pm$.028	.369 $\pm$.053	.821 $\pm$.016
not-MIWAE	.587 $\pm$.027	.540 $\pm$.043	.839 $\pm$.013	.558 $\pm$.027	.394 $\pm$.049	.818 $\pm$.017
MIRACLE	.597 $\pm$.027	.473 $\pm$.046	.851 $\pm$.014	.564 $\pm$.027	.264 $\pm$.051	.825 $\pm$.016
ReMasker	.592 $\pm$.027	.480 $\pm$.046	.817 $\pm$.019	.566 $\pm$.028	.279 $\pm$.050	.782 $\pm$.021
DrIM _{BASE}	.640 $\pm$.024	.617 $\pm$.041	.894 $\pm$.010	.625 $\pm$.024	.496 $\pm$.043	.886 $\pm$.009
DrIM _{FINE}	.659 $\pm$.025	.658 $\pm$.032	.905 $\pm$.007	.653 $\pm$.025	.553 $\pm$.040	.915 $\pm$.006

evaluate the performance of post-imputation prediction. Specifically, we use three metrics detailed in Hansen et al. (2023): classification performance (F_1 score), model selection performance (Model), and feature selection performance (Feature). For a comprehensive evaluation procedure, please refer to Appendix A.6.

4.2 RESULTS

Imputation fidelity. Figure 2 shows the imputation fidelity of DrIM and baseline models at a 0.3 missingness rate. DrIM consistently demonstrates competitive performance across the four missingness mechanisms (MCAR, MAR, MNARQ, and MNARL), with consistently low ARSMAPE scores. This indicates that our model is robust to different missing data mechanisms, which are often encountered in real-world datasets. Notably, DrIM_{FINE} benefits from the contrastive learning strategy, as shown by its superior performance compared to DrIM_{BASE}. This improvement suggests that the contrastive learning approach enhances DrIM’s ability to identify relevant nearest neighbors, resulting in more accurate imputation. These enhanced imputations are crucial for post-imputation tasks, such as MLu, that depend on the quality of the imputed data.

Machine learning utility. We evaluate the post-imputation performance of DrIM and baseline models in terms of MLu. As shown in Table 1, DrIM consistently achieves the highest metric scores in F_1 score, model selection, and feature selection. Notably, DrIM_{BASE} outperforms other baselines, such as GAIN and VAEAC, which require training, even though these models achieve similar performance to DrIM. This highlights the superior performance of DrIM_{BASE} without the need for training. This suggests that our approach is cost-effective, providing strong performance without the need for extensive training. Additionally, we confirmed that kNNI demonstrates competitive performance compared to other baseline models. Given that our method outperforms kNNI, this empirically vali-

Table 2: **Effect of contrastive learning** under MAR. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. Values in parentheses indicate the performance difference compared to $\text{DrIM}_{\text{BASE}}$, and the red highlights the positive improvement.

Rate	model	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
0.2	$\text{DrIM}_{\text{BASE}}$	.673 $\pm$.024	.732 $\pm$.029	.934 $\pm$.006
	$\text{DrIM}_{\text{FINE}}$	.684 $\pm$.025 (+1.64%)	.773 $\pm$.023 (+5.59%)	.948 $\pm$.004 (+1.50%)
0.4	$\text{DrIM}_{\text{BASE}}$	.607 $\pm$.024	.452 $\pm$.044	.837 $\pm$.015
	$\text{DrIM}_{\text{FINE}}$	.639 $\pm$.025 (+5.27%)	.556 $\pm$.035 (+23.03%)	.853 $\pm$.012 (+1.91%)
0.6	$\text{DrIM}_{\text{BASE}}$	.545 $\pm$.024	.231 $\pm$.052	.684 $\pm$.024
	$\text{DrIM}_{\text{FINE}}$	.589 $\pm$.024 (+8.08%)	.373 $\pm$.040 (+61.47%)	.734 $\pm$.026 (+7.32%)
0.8	$\text{DrIM}_{\text{BASE}}$	.477 $\pm$.023	.072 $\pm$.054	.446 $\pm$.041
	$\text{DrIM}_{\text{FINE}}$	.534 $\pm$.024 (+11.96%)	.257 $\pm$.047 (+257.64%)	.604 $\pm$.035 (+35.62%)

dates that the distance between representations of BERT is effective. In Appendix A.9.2, we provide more comprehensive results under different missingness mechanisms and missingness rates. Additionally, you can find the MLu performance for each dataset.

Effect of contrastive learning. The results presented in Table 2 indicate that incorporating the contrastive learning framework into our proposed method improves performance. As shown in Table 1 and Table 2, $\text{DrIM}_{\text{FINE}}$ outperforms $\text{DrIM}_{\text{BASE}}$, supporting our claim in Section 3.4 that fine-tuning with contrastive learning allows the representations of BERT to be trained in a way that performs metric learning to maximize the mutual information between input sequences. Furthermore, we observe performance improvements not only in MAR but also in MNARL, further validating the effectiveness of our approach. Additional results for different missing data mechanisms (MCAR, MNARL, MNARQ) can be found in Appendix A.9.4.

Sensitivity analysis. To more comprehensively evaluate the post-imputation performance of DrIM , we also conducted a sensitivity analysis by varying the missingness rate of the datasets. Figure 3 shows that although the performance of each model decreases as the missingness rate increases, both $\text{DrIM}_{\text{BASE}}$ and $\text{DrIM}_{\text{FINE}}$ exhibit stable performance compared to the baseline models. This reflects the effectiveness of leveraging the language model’s representation learning capabilities, even without fine-tuning. Notably, $\text{DrIM}_{\text{FINE}}$ consistently outperforms the baseline models in terms of F_1 score and feature selection performance. Although the model selection performance of DrIM crosses with VAEAC as the missingness rate changes from 0.4 to 0.6, $\text{DrIM}_{\text{FINE}}$ still demonstrates superior performance compared to other baseline models. This indicates that fine-tuning with contrastive learning continues to be effective even at higher missingness rates. Additional results for different missing data mechanisms can be found in Appendix A.9.5.

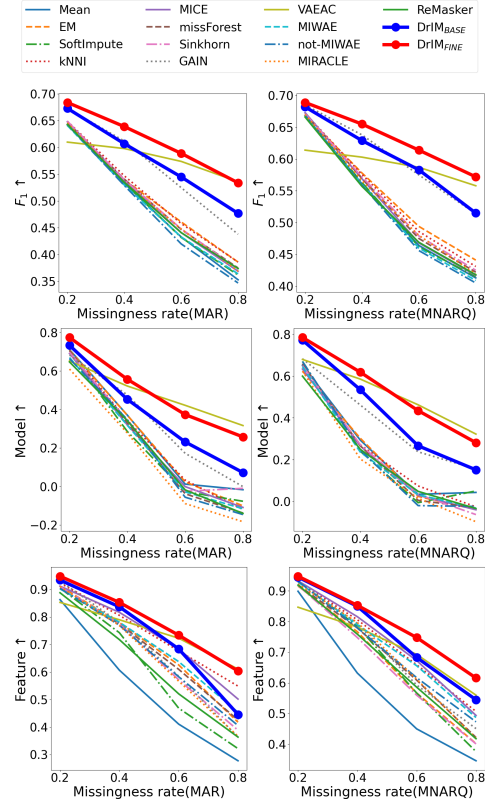


Figure 3: **Sensitivity analysis for missingness rates.** The results MAR and MNARQ are shown, with the first-row representing classification performance (F_1 score) and the last two rows displaying model selection and feature selection performance. The means of the average across 10 datasets and 10 repeated experiments are reported.

5 CONCLUSIONS

This paper introduces DrIM, a context-driven nearest-neighbor-based imputation method specifically designed to handle incomplete heterogeneous tabular datasets. Our proposed method leverages the representation learning capabilities of language models and integrates a contrastive learning framework for fine-tuning. In other words, in our proposed method, BERT performs dimensionality reduction on (numerical) tabular datasets in text format through the bidirectional attention mechanism, and we perform metric learning to maximize mutual information between observed and missing data entries. Our extensive experiments across 10 real-world datasets demonstrate that DrIM consistently outperforms existing imputation methods, achieving state-of-the-art post-imputation performance in terms of MLu across various missing data mechanisms, including MCAR, MAR, MNARL, and MNARQ.

In this paper we have not explicitly demonstrated how the neighbors using a BERT model for kNNI specifically improves downstream task performance. However, we believe our proposed method provides insights into the relationship between learned metrics obtained through contrastive learning and conditional independence among variables by emphasizing the enhancement of machine learning utility rather than statistical fidelity. This connection between contextual embeddings derived from BERT and downstream task performance may serve as a critical clue for further exploration, which we leave for future research.

REPRODUCIBILITY STATEMENT

Detailed information about the datasets we used can be found in Appendix A.4. For the reproducibility of the baseline models and our proposed model, detailed experimental settings are provided in Appendix A.5. The reproduced codes for the baseline models and our proposed model are available in the supplementary material. For a comprehensive evaluation of MLu, please refer to Appendix A.6 for the detailed evaluation procedure and ML model configuration.

REFERENCES

- Loai AbedAllah and Ilan Shimshoni. K-means over incomplete datasets using mean euclidean distance. In *IAPR International Conference on Machine Learning and Data Mining in Pattern Recognition*, 2016. URL <https://api.semanticscholar.org/CorpusID:6796762>.
- David W. Aha, Dennis F. Kibler, and Marc K. Albert. Instance-based learning algorithms. *Machine Learning*, 6:37–66, 1991. URL <https://api.semanticscholar.org/CorpusID:207670313>.
- Naomi S Altman. An introduction to kernel and nearest-neighbor nonparametric regression. *The American Statistician*, 46(3):175–185, 1992.
- Seunghwan An and Jong-June Jeon. Distributional learning of variational autoencoder: Application to synthetic data generation. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023. URL <https://openreview.net/forum?id=GxL6PrmEUw>.
- Dhanya Pramod Anil Jadhav and Krishnan Ramanathan. Comparison of performance of data imputation methods for numeric dataset. *Applied Artificial Intelligence*, 33(10):913–933, 2019. doi: 10.1080/08839514.2019.1637138. URL <https://doi.org/10.1080/08839514.2019.1637138>.
- Philip Bachman, R Devon Hjelm, and William Buchwalter. Learning representations by maximizing mutual information across views. *Advances in neural information processing systems*, 32, 2019.
- Mohamed Ishmael Belghazi, Aristide Baratin, Sai Rajeshwar, Sherjil Ozair, Yoshua Bengio, Aaron Courville, and Devon Hjelm. Mutual information neural estimation. In *International conference on machine learning*, pp. 531–540. PMLR, 2018.
- Sid Black, Leo Gao, Phil Wang, Connor Leahy, and Stella Biderman. Gpt-neo: Large scale autoregressive language modeling with mesh-tensorflow. 2021. URL <https://api.semanticscholar.org/CorpusID:245758737>.

- Vadim Borisov, Kathrin Sessler, Tobias Leemann, Martin Pawelczyk, and Gjergji Kasneci. Language models are realistic tabular data generators. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=cEYgmQNOeI>.
- Leo Breiman. Random forests. *Machine learning*, 45:5–32, 2001.
- Tianqi Chen and Carlos Guestrin. Xgboost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '16, pp. 785–794, New York, NY, USA, 2016. Association for Computing Machinery. ISBN 9781450342322. doi: 10.1145/2939672.2939785. URL <https://doi.org/10.1145/2939672.2939785>.
- Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. A simple framework for contrastive learning of visual representations. In *International conference on machine learning*, pp. 1597–1607. PMLR, 2020.
- Juan Colonna, Eduardo Nakamura, Marco Cristo, and Marcelo Gordo. Anuran Calls (MFCCs). UCI Machine Learning Repository, 2015. DOI: <https://doi.org/10.24432/C5CC9H>.
- Paulo Cortez, A. Cerdeira, F. Almeida, T. Matos, and J. Reis. Wine Quality. UCI Machine Learning Repository, 2009. DOI: <https://doi.org/10.24432/C56S3T>.
- David R Cox. The regression analysis of binary sequences. *Journal of the Royal Statistical Society: Series B (Methodological)*, 20(2):215–232, 1958.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. In *North American Chapter of the Association for Computational Linguistics*, 2019.
- Zhicheng Ding, Jiahao Tian, Zhenkai Wang, Jinman Zhao, and Siyang Li. Data imputation using large language model to accelerate recommendation system. 2024. URL <https://api.semanticscholar.org/CorpusID:271213203>.
- Tianyu Du, Luca Melis, and Ting Wang. Remasker: Imputing tabular data with masked autoencoding. In *The Twelfth International Conference on Learning Representations*, 2024a. URL <https://openreview.net/forum?id=KI9NqjLVDT>.
- Wenyou Du, Yichen Wang, Guanglei Meng, and Yuming Guo. Privacy-preserving vertical federated knn feature imputation method. *Electronics*, 2024b. URL <https://api.semanticscholar.org/CorpusID:267038107>.
- Alireza Farhangfar, Lukasz A Kurgan, and Witold Pedrycz. A novel framework for imputation of missing values in databases. *IEEE Transactions on Systems, Man, and Cybernetics-Part A: Systems and Humans*, 37(5):692–709, 2007.
- Pedro J. García-Laencina, José-Luis Sancho-Gómez, and Aníbal R. Figueiras-Vidal. Pattern classification with missing data: a review. *Neural Computing and Applications*, 19:263–282, 2010. URL <https://api.semanticscholar.org/CorpusID:3351246>.
- Ian J. Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative adversarial nets. In *Proceedings of the 27th International Conference on Neural Information Processing Systems - Volume 2*, NIPS'14, pp. 2672–2680, Cambridge, MA, USA, 2014. MIT Press.
- Lasse Hansen, Nabeel Seedat, Mihaela van der Schaar, and Andrija Petrovic. Reimagining synthetic tabular data generation through data-centric AI: A comprehensive benchmark. In *Thirty-seventh Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2023. URL <https://openreview.net/forum?id=dK1Rs1o0Ij>.
- Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. In H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 6840–6851. Curran Associates, Inc., 2020. URL https://proceedings.neurips.cc/paper_files/paper/2020/file/4c5bcfec8584af0d967f1ab10179ca4b-Paper.pdf.

- Tin Kam Ho. The random subspace method for constructing decision forests. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 20(8):832–844, 1998. doi: 10.1109/34.709601.
- Xin Huang, Ashish Khetan, Milan Cvitkovic, and Zohar Karnin. Tabtransformer: Tabular data modeling using contextual embeddings, 2020.
- Niels Bruun Ipsen, Pierre-Alexandre Mattei, and Jes Frellsen. not- $\{miwae\}$: Deep generative modelling with missing not at random data. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=tu29GQT0JFy>.
- Oleg Ivanov, Michael Figurnov, and Dmitry Vetrov. Variational autoencoder with arbitrary conditioning. In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=SyxtJh0qYm>.
- Daniel Jarrett, Bogdan C Cebere, Tennison Liu, Alicia Curth, and Mihaela van der Schaar. Hyper-Impute: Generalized iterative imputation with automatic model selection. In Kamalika Chaudhuri, Stefanie Jegelka, Le Song, Csaba Szepesvari, Gang Niu, and Sivan Sabato (eds.), *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pp. 9916–9937. PMLR, 17–23 Jul 2022. URL <https://proceedings.mlr.press/v162/jarrett22a.html>.
- José M. Jerez, Ignacio Molina, Pedro J. García-Laencina, Emilio Alba, Nuria Ribelles, Miguel Martín, and Leonardo Franco. Missing data imputation using statistical and machine learning methods in a real breast cancer problem. *Artificial Intelligence in Medicine*, 50(2):105–115, 2010. ISSN 0933-3657. doi: <https://doi.org/10.1016/j.artmed.2010.05.002>. URL <https://www.sciencedirect.com/science/article/pii/S0933365710000679>.
- Chao Jiang and Zijiang Yang. Cknni: An improved knn-based missing value handling technique. In De-Shuang Huang and Kyungsook Han (eds.), *Advanced Intelligent Computing Theories and Applications*, pp. 441–452, Cham, 2015. Springer International Publishing. ISBN 978-3-319-22053-6.
- Yuxin Jiang and Wei Wang. Deep continuous prompt for contrastive learning of sentence embeddings, 2022.
- Martti Juhola and Jorma Laurikkala. On metricity of two heterogeneous measures in the presence of missing values. *Artificial Intelligence Review*, 28:163–178, 2007. URL <https://api.semanticscholar.org/CorpusID:33847594>.
- Guolin Ke, Qi Meng, Thomas Finley, Taifeng Wang, Wei Chen, Weidong Ma, Qiwei Ye, and Tie-Yan Liu. Lightgbm: A highly efficient gradient boosting decision tree. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL https://proceedings.neurips.cc/paper_files/paper/2017/file/6449f44a102fde848669bdd9eb6b76fa-Paper.pdf.
- Phimmarin Keerin and Tossapon Boongoen. Improved knn imputation for missing values in gene expression data. *Computers, Materials & Continua*, 2022. URL <https://api.semanticscholar.org/CorpusID:239131676>.
- Diederik P. Kingma and Max Welling. Auto-Encoding Variational Bayes. In *2nd International Conference on Learning Representations, ICLR 2014, Banff, AB, Canada, April 14-16, 2014, Conference Track Proceedings*, 2014.
- Trent Kyono, Yao Zhang, Alexis Bellot, and Mihaela van der Schaar. MIRACLE: Causally-aware imputation via learning missing data mechanisms. In A. Beygelzimer, Y. Dauphin, P. Liang, and J. Wortman Vaughan (eds.), *Advances in Neural Information Processing Systems*, 2021. URL <https://openreview.net/forum?id=GzeqcAUFGl0>.
- Roderick JA Little and Donald B Rubin. *Statistical analysis with missing data*, volume 793. John Wiley & Sons, 2019.

- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. Roberta: A robustly optimized BERT pretraining approach. *CoRR*, abs/1907.11692, 2019. URL <http://arxiv.org/abs/1907.11692>.
- Volker Lohweg. Banknote Authentication. UCI Machine Learning Repository, 2012. DOI: <https://doi.org/10.24432/C55P57>.
- Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *International Conference on Learning Representations*, 2017. URL <https://api.semanticscholar.org/CorpusID:53592270>.
- Pierre-Alexandre Mattei and Jes Frellsen. Miwae: Deep generative modelling and imputation of incomplete data sets. In *International Conference on Machine Learning*, 2019.
- Rahul Mazumder, Trevor Hastie, and Robert Tibshirani. Spectral regularization algorithms for learning large incomplete matrices. *The Journal of Machine Learning Research*, 11:2287–2322, 2010.
- Yinan Mei, Shaoxu Song, Chenguang Fang, Haifeng Yang, Jingyun Fang, and Jiang Long. Capturing semantics for imputation with pre-trained language models. In *2021 IEEE 37th International Conference on Data Engineering (ICDE)*, pp. 61–72, 2021. doi: 10.1109/ICDE51399.2021.00013.
- Deshui Miao, Jiaqi Zhang, Wenbo Xie, Jian Song, Xin Li, Lijuan Jia, and Ning Guo. Simple contrastive representation adversarial learning for nlp tasks. *arXiv preprint arXiv:2111.13301*, 2021.
- Xiaoye Miao, Yangyang Wu, Lu Chen, Yunjun Gao, and Jianwei Yin. An experimental survey of missing data imputation algorithms. *IEEE Transactions on Knowledge and Data Engineering*, 2022.
- Lluís A. Belanche Muñoz and Jerónimo Hernández-González. Similarity networks for heterogeneous data. In *The European Symposium on Artificial Neural Networks*, 2012. URL <https://api.semanticscholar.org/CorpusID:6885211>.
- Boris Muzellec, Julie Josse, Claire Boyer, and Marco Cuturi. Missing data imputation using optimal transport. In *International Conference on Machine Learning*, pp. 7130–7140. PMLR, 2020.
- Warwick Nash, Tracy Sellers, Simon Talbot, Andrew Cawthorn, and Wes Ford. Abalone. UCI Machine Learning Repository, 1994. DOI: <https://doi.org/10.24432/C55C7W>.
- Alfredo Nazábal, Pablo M. Olmos, Zoubin Ghahramani, and Isabel Valera. Handling incomplete heterogeneous data using vaes. *Pattern Recognition*, 107:107501, 2020. ISSN 0031-3203. doi: <https://doi.org/10.1016/j.patcog.2020.107501>. URL <https://www.sciencedirect.com/science/article/pii/S0031320320303046>.
- Anam Nazir, Muhammad Nadeem Cheema, and Ze Wang. Chatgpt-based biological and psychological data imputation. *Meta-Radiology*, 1(3):100034, 2023. ISSN 2950-1628. doi: <https://doi.org/10.1016/j.metrad.2023.100034>. URL <https://www.sciencedirect.com/science/article/pii/S2950162823000346>.
- Fulufhelo V Nelwamondo, Shakir Mohamed, and Tshilidzi Marwala. Missing data: A comparison of neural network and expectation maximization techniques. *Current Science*, pp. 1514–1521, 2007.
- Aaron van den Oord, Yazhe Li, and Oriol Vinyals. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748*, 2018.
- Alec Radford and Karthik Narasimhan. Improving language understanding by generative pre-training. 2018. URL <https://api.semanticscholar.org/CorpusID:49313245>.
- Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. Language models are unsupervised multitask learners. 2019. URL <https://api.semanticscholar.org/CorpusID:160025533>.

- Donald B. Rubin and Nathaniel Schenker. Multiple imputation for interval estimation from simple random samples with ignorable nonresponse. *Journal of the American Statistical Association*, 81: 366–374, 1986.
- Avraham Ruderman, Mark Reid, Darío García-García, and James Petterson. Tighter variational representations of f-divergences via restriction to probability measures. *arXiv preprint arXiv:1206.4664*, 2012.
- Miriam Seoane Santos, Jastin Pompeu Soares, Pedro Henriques Abreu, Hélder Araújo, and João Santos. Influence of data distribution in missing data imputation. In *Artificial Intelligence in Medicine: 16th Conference on Artificial Intelligence in Medicine, AIME 2017, Vienna, Austria, June 21–24, 2017, Proceedings 16*, pp. 285–294. Springer, 2017.
- Miriam Seoane Santos, Pedro Henriques Abreu, Szymon Wilk, and João A. M. Santos. How distance metrics influence missing data imputation with k-nearest neighbours. *Pattern Recognit. Lett.*, 136:111–119, 2020. URL <https://api.semanticscholar.org/CorpusID:219756021>.
- Hiroaki Sasaki and Takashi Takenouchi. Representation learning for maximization of mi, nonlinear ica and nonlinear subspaces with robust density ratio estimation. *Journal of Machine Learning Research*, 23(231):1–55, 2022.
- David Slate. Letter Recognition. UCI Machine Learning Repository, 1991. DOI: <https://doi.org/10.24432/C5ZP40>.
- Gowthami Somepalli, Micah Goldblum, Avi Schwarzschild, C. Bayan Bruss, and Tom Goldstein. Saint: Improved neural networks for tabular data via row attention and contrastive pre-training. *ArXiv*, abs/2106.01342, 2021. URL <https://api.semanticscholar.org/CorpusID:235293989>.
- Craig Stanfill and David L. Waltz. Toward memory-based reasoning. *Commun. ACM*, 29:1213–1228, 1986. URL <https://api.semanticscholar.org/CorpusID:16624499>.
- Daniel J Stekhoven and Peter Bühlmann. Missforest—non-parametric missing value imputation for mixed-type data. *Bioinformatics*, 28(1):112–118, 2012.
- Masashi Sugiyama, Taiji Suzuki, and Takafumi Kanamori. *Density ratio estimation in machine learning*. Cambridge University Press, 2012.
- Huachun Tan, Guangdong Feng, Jianshuai Feng, Wuhong Wang, Yu-Jin Zhang, and Feng Li. A tensor-based method for missing traffic data completion. *Transportation Research Part C: Emerging Technologies*, 28:15–27, 2013. ISSN 0968-090X. doi: <https://doi.org/10.1016/j.trc.2012.12.007>. URL <https://www.sciencedirect.com/science/article/pii/S0968090X12001532>. Euro Transportation: selected paper from the EWGT Meeting, Padova, September 2009.
- Jian Tang, Meng Qu, Mingzhe Wang, Ming Zhang, Jun Yan, and Qiaozhu Mei. Line: Large-scale information network embedding. *Proceedings of the 24th International Conference on World Wide Web*, 2015. URL <https://api.semanticscholar.org/CorpusID:8399404>.
- Agostino Tarsitano and Marianna Falcone. Missing-values adjustment for mixed-type data. *Journal of Probability and Statistics*, 2011:1–20, 2011. URL <https://api.semanticscholar.org/CorpusID:46672653>.
- Tressy Thomas and Enayat Rajabi. Addressing missing data in a healthcare dataset using an improved knn algorithm. In Maciej Paszynski, Dieter Kranzlmüller, Valeria V. Krzhizhanovskaya, Jack J. Dongarra, and Peter M.A. Sloot (eds.), *Computational Science – ICCS 2021*, pp. 223–230, Cham, 2021. Springer International Publishing. ISBN 978-3-030-77977-1.
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, Aurelien Rodriguez, Armand Joulin, Edouard Grave, and Guillaume Lample. Llama: Open and efficient foundation language models, 2023. URL <https://arxiv.org/abs/2302.13971>.

- Olga G. Troyanskaya, Michael N. Cantor, Gavin Sherlock, Patrick O. Brown, Trevor J. Hastie, Robert Tibshirani, David Botstein, and Russ B. Altman. Missing value estimation methods for dna microarrays. *Bioinformatics*, 17 6:520–5, 2001. URL <https://api.semanticscholar.org/CorpusID:1105917>.
- Stef van Buuren. Flexible imputation of missing data. 2012.
- Stef van Buuren and Karin Groothuis-Oudshoorn. mice: Multivariate imputation by chained equations in r. *Journal of Statistical Software*, 45(3):1–67, 2011. doi: 10.18637/jss.v045.i03. URL <https://www.jstatsoft.org/index.php/jss/article/view/v045i03>.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf.
- Yizhu Wen, Kai Yi, Jing Ke, and Yiqing Shen. Diffimpute: Tabular data imputation with denoising diffusion probabilistic model. *ArXiv*, abs/2403.13863, 2024. URL <https://api.semanticscholar.org/CorpusID:268554183>.
- D Randall Wilson and Tony R Martinez. Improved heterogeneous distance functions. *Journal of artificial intelligence research*, 6:1–34, 1997.
- William Wolberg, Olvi Mangasarian, Nick Street, , and W. Street. Breast Cancer Wisconsin (Diagnostic). UCI Machine Learning Repository, 1993. DOI: <https://doi.org/10.24432/C5DW2B>.
- Xindong Wu, Vipin Kumar, J Ross Quinlan, Joydeep Ghosh, Qiang Yang, Hiroshi Motoda, Geoffrey J McLachlan, Angus Ng, Bing Liu, S Yu Philip, et al. Top 10 algorithms in data mining. *Knowledge and information systems*, 14(1):1–37, 2008.
- Xing Wu, Chaochen Gao, Liangjun Zang, Jizhong Han, Zhongyuan Wang, and Songlin Hu. ES-imCSE: Enhanced sample building method for contrastive learning of unsupervised sentence embedding. In Nicoletta Calzolari, Chu-Ren Huang, Hansaem Kim, James Pustejovsky, Leo Wanner, Key-Sun Choi, Pum-Mo Ryu, Hsin-Hsi Chen, Lucia Donatelli, Heng Ji, Sadao Kurohashi, Patrizia Paggio, Nianwen Xue, Seokhwan Kim, Younggyun Hahm, Zhong He, Tony Kyungil Lee, Enrico Santus, Francis Bond, and Seung-Hoon Na (eds.), *Proceedings of the 29th International Conference on Computational Linguistics*, pp. 3898–3907, Gyeongju, Republic of Korea, October 2022. International Committee on Computational Linguistics. URL <https://aclanthology.org/2022.coling-1.342>.
- I-Cheng Yeh. Concrete Compressive Strength. UCI Machine Learning Repository, 1998. DOI: <https://doi.org/10.24432/C5PK67>.
- Pengcheng Yin, Graham Neubig, Wen-tau Yih, and Sebastian Riedel. TaBERT: Pretraining for joint understanding of textual and tabular data. In Dan Jurafsky, Joyce Chai, Natalie Schluter, and Joel Tetreault (eds.), *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pp. 8413–8426, Online, July 2020. Association for Computational Linguistics. doi: 10.18653/v1/2020.acl-main.745. URL <https://aclanthology.org/2020.acl-main.745>.
- Jinsung Yoon, James Jordon, and Mihaela van der Schaar. GAIN: Missing data imputation using generative adversarial nets. In Jennifer Dy and Andreas Krause (eds.), *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pp. 5689–5698. PMLR, 10–15 Jul 2018. URL <https://proceedings.mlr.press/v80/yoon18a.html>.
- He Zhao, Ke Sun, Amir Dezfouli, and Edwin V. Bonilla. Transformed distribution matching for missing value imputation. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett (eds.), *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pp. 42159–42186. PMLR, 23–29 Jul 2023. URL <https://proceedings.mlr.press/v202/zhao23h.html>.

810 Encarnación Álvarez Verdejo, Pablo J. Moya-Fernández, and Juan F. Muñoz-Rosas. Single impu-
811 tation methods and confidence intervals for the gini index. *Mathematics*, 9(24), 2021. ISSN
812 2227-7390. doi: 10.3390/math9243252. URL [https://www.mdpi.com/2227-7390/9/](https://www.mdpi.com/2227-7390/9/24/3252)
813 [24/3252](https://www.mdpi.com/2227-7390/9/24/3252).
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863

A APPENDIX

A.1 PROOF OF THEOREM 1

Proof. Sugiyama et al. (2012); Sasaki & Takenouchi (2022) showed that the minimizer of $J_{LR}(\theta, \eta)$ is given by

$$r(\mathbf{t}, \mathbf{u}; \theta, \eta) = \log \frac{p(\mathbf{t}, \mathbf{u})}{p(\mathbf{t})p(\mathbf{u})}.$$

Then, for the minimizer of $J_{LR}(\theta, \eta)$, we have

$$\begin{aligned} \left| r(\mathbf{t}, \mathbf{e}; \theta, \eta) - r(\mathbf{t}', \mathbf{e}; \theta, \eta) \right| &= \left| \log \frac{p(\mathbf{t}, \mathbf{e})}{p(\mathbf{t})p(\mathbf{e})} - \log \frac{p(\mathbf{t}', \mathbf{e})}{p(\mathbf{t}')p(\mathbf{e})} \right| \\ &= \left| \log \frac{p(\mathbf{e} | \mathbf{t})}{p(\mathbf{e} | \mathbf{t}')} \right|, \end{aligned}$$

and

$$\begin{aligned} \left| r(\mathbf{t}, \mathbf{e}; \theta, \eta) - r(\mathbf{t}', \mathbf{e}; \theta, \eta) \right| &= \left| \Phi(g(\mathbf{t}, \mathbf{e}; \theta); \eta) - \Phi(g(\mathbf{t}', \mathbf{e}; \theta); \eta) \right| \\ &= \left| \eta^\top (g(\mathbf{t}, \mathbf{e}; \theta) - g(\mathbf{t}', \mathbf{e}; \theta)) \right| \quad (\text{the scoring function is a linear map}) \\ &\leq \|\eta\| \cdot \|g(\mathbf{t}, \mathbf{e}; \theta) - g(\mathbf{t}', \mathbf{e}; \theta)\|, \end{aligned}$$

where the last inequality follows from Cauchy-Schwarz inequality.

The proof is complete. $\square$

A.2 PROOF OF THEOREM 2

Proof. Ruderman et al. (2012); Belghazi et al. (2018) showed that

$$\begin{aligned} I((\mathbf{t}, \mathbf{e}), (\mathbf{u}, \mathbf{e})) &\geq J_{CL}(\theta) \\ &= \mathbb{E}_{q(\mathbf{t}, \mathbf{u}, \mathbf{e})} [\cos(g(\mathbf{t}, \mathbf{e}; \theta), g(\mathbf{u}, \mathbf{e}; \theta))] - \log \mathbb{E}_{p(\mathbf{t}, \mathbf{e})p(\mathbf{u}, \mathbf{e})} \left[\exp \left(\cos(g(\mathbf{t}, \mathbf{e}; \theta), g(\mathbf{u}, \mathbf{e}; \theta)) \right) \right], \end{aligned}$$

where $I((\mathbf{t}, \mathbf{e}), (\mathbf{u}, \mathbf{e}))$ is the mutual information between $\mathbf{t}$ and $\mathbf{u}$.

Therefore, if we obtain θ such that the maximizer of $J_{CL}(\theta)$, $g(\mathbf{t}, \mathbf{e}; \theta)$ and $g(\mathbf{u}, \mathbf{e}; \theta)$ provide the representation vectors that maximizes the lower bound of the mutual information between $\mathbf{t}$ and $\mathbf{u}$.

The proof is complete. $\square$

A.3 PROOF OF THEOREM 3

The proof is a restatement of the results from the paper by Oord et al. (2018).

Proof. $\mathcal{L}(\theta)$ is the categorical cross-entropy for classifying the positive sample correctly, with the following expression for the prediction from BERT:

$$\frac{\exp \left(\cos(g(\mathbf{t}^{(i)}, \mathbf{e}; \theta), g(\mathbf{t}^{(i+)}, \mathbf{e}; \theta)) \right)}{\exp \left(\cos(g(\mathbf{t}^{(i)}, \mathbf{e}; \theta), g(\mathbf{t}^{(i+)}, \mathbf{e}; \theta)) \right) + \sum_{l=1, l \neq i}^B \exp \left(\cos(g(\mathbf{t}^{(i)}, \mathbf{e}; \theta), g(\mathbf{t}^{(l)}, \mathbf{e}; \theta)) \right)}.$$

With respect to the i th sample, the optimal probability for $\mathcal{L}(\theta)$ is denoted as

$$\Pr(i+ \text{ is the positive sample} \mid \mathbf{t}^{(l)}, l \in \{i+, 1, 2, \dots, B\}).$$

Note that the positive sample $\mathbf{t}^{(i+)}$ is drawn from the conditional distribution of $q(\mathbf{u} \mid \mathbf{t}^{(i)})$ and $\Pr(l \text{ is the positive sample}) = 1/B$ for all $l \in \{i+, 1, 2, \dots, B\}$ and $l \neq i$. Then, we have

$$\begin{aligned}
& \Pr(i+ \text{ is the positive sample} \mid \mathbf{t}^{(l)}, l \in \{i+, 1, 2, \dots, B\}) \\
&= \frac{p(\mathbf{t}^{(l)}, l \in \{i+, 1, 2, \dots, B\} \mid i+ \text{ is the positive sample}) \Pr(i+ \text{ is the positive sample})}{\sum_{j \in \{i+, 1, 2, \dots, B\}, j \neq i} p(\mathbf{t}^{(l)}, l \in \{i+, 1, 2, \dots, B\} \mid j \text{ is the positive sample}) \Pr(j \text{ is the positive sample})} \\
&= \frac{q(\mathbf{t}^{(i+)} \mid \mathbf{t}^{(i)}) \prod_{l=1, l \neq i}^B p(\mathbf{t}^{(l)}, \mathbf{e})}{q(\mathbf{t}^{(i+)} \mid \mathbf{t}^{(i)}) \prod_{l=1, l \neq i}^B p(\mathbf{t}^{(l)}, \mathbf{e}) + \sum_{j=1, j \neq i}^B q(\mathbf{t}^{(j)} \mid \mathbf{t}^{(i)}) \prod_{l \in \{i+, 1, 2, \dots, B\}, l \neq j, l \neq i} p(\mathbf{t}^{(l)}, \mathbf{e})} \\
&= \frac{\frac{q(\mathbf{t}^{(i+)} \mid \mathbf{t}^{(i)})}{p(\mathbf{t}^{(i+)}, \mathbf{e})}}{\frac{q(\mathbf{t}^{(i+)} \mid \mathbf{t}^{(i)})}{p(\mathbf{t}^{(i+)}, \mathbf{e})} + \sum_{j=1, j \neq i}^B \frac{q(\mathbf{t}^{(j)} \mid \mathbf{t}^{(i)})}{p(\mathbf{t}^{(j)}, \mathbf{e})}} \\
&= \frac{\frac{p(\mathbf{t}^{(i)}, \mathbf{t}^{(i+)}, \mathbf{e})}{p(\mathbf{t}^{(i)}, \mathbf{e})p(\mathbf{t}^{(i+)}, \mathbf{e})}}{\frac{p(\mathbf{t}^{(i)}, \mathbf{t}^{(i+)}, \mathbf{e})}{p(\mathbf{t}^{(i)}, \mathbf{e})p(\mathbf{t}^{(i+)}, \mathbf{e})} + \sum_{j=1, j \neq i}^B \frac{p(\mathbf{t}^{(i)}, \mathbf{t}^{(j)}, \mathbf{e})}{p(\mathbf{t}^{(i)}, \mathbf{e})q(\mathbf{t}^{(j)}, \mathbf{e})}}.
\end{aligned}$$

Therefore, for the optimal value of the prediction from BERT with θ such that the maximizer of $\mathcal{L}(\theta)$,

$$\cos(g(\mathbf{t}^{(i)}, \mathbf{e}; \theta), g(\mathbf{t}^{(i+)}, \mathbf{e}; \theta)) \propto \log \frac{p(\mathbf{t}^{(i)}, \mathbf{t}^{(i+)}, \mathbf{e})}{p(\mathbf{t}^{(i)}, \mathbf{e})p(\mathbf{t}^{(i+)}, \mathbf{e})} = \log \frac{p(\mathbf{t}^{(i)} \mid \mathbf{t}^{(i+)})}{p(\mathbf{t}^{(i)}, \mathbf{e})},$$

where the (unnormalized) density ratio is modeled as log-linear (Bachman et al., 2019), and $q(\mathbf{t}^{(i)}, \mathbf{t}^{(i+)})$, $q(\mathbf{t}^{(i)})$, and $q(\mathbf{t}^{(i+)})$ denote the joint and marginal distributions of $\mathbf{t}^{(i)}$ and $\mathbf{t}^{(i+)}$, respectively. $\square$

A.4 DATASET DESCRIPTIONS

Download links.

- abalone (Nash et al., 1994):
<https://archive.ics.uci.edu/dataset/1/abalone>
- anuran (Colonna et al., 2015):
<https://archive.ics.uci.edu/dataset/406/anuran+calls+mfccs>
- banknote (Lohweg, 2012):
<https://archive.ics.uci.edu/dataset/267/banknote+authentication>
- breast (Wolberg et al., 1993):
<https://archive.ics.uci.edu/dataset/17/breast+cancer+wisconsin+diagnostic>
- concrete (Yeh, 1998):
<https://archive.ics.uci.edu/dataset/165/concrete+compressive+strength>
- kings (CC0: Public Domain):
<https://www.kaggle.com/datasets/harlfoxem/housesalesprediction>
- letter (Slate, 1991):
<https://archive.ics.uci.edu/dataset/59/letter+recognition>
- loan (CC0: Public Domain):
<https://www.kaggle.com/datasets/teertha/personal-loan-modeling>
- redwine (Cortez et al., 2009):
<https://archive.ics.uci.edu/dataset/186/wine+quality>
- whitewine (Cortez et al., 2009):
<https://archive.ics.uci.edu/dataset/186/wine+quality>

Table 3: **Description of datasets.** #C represents the number of continuous variables. #D denotes the number of categorical (discrete) variables. The ‘Target’ refers to the variable used as the response variable in a classification task to evaluate machine learning utility.

Dataset	Split	#C	#D	Target
abalone	3.3K/0.8K	7	2	Rings
anuran	5.7K/1.5K	22	3	Species
banknote	1.1K/0.3K	4	1	class
breast	0.5K/0.1K	30	1	Diagnosis
concrete	0.8K/0.2K	8	1	Age
kings	17.3K/4.3K	11	7	grade
letter	16K/4K	16	1	lettr
loan	4K/1K	5	6	Personal Loan
redwine	1.3K/0.3K	11	1	quality
whitewine	3.9K/1K	11	1	quality

Algorithm 1 Imputation procedure of DrIM**Input:** Incomplete observation: $(\mathbf{x}, \mathbf{m})$, Columns name: $C = \{C_1, C_2, \dots, C_p\}$,Pre-trained BERT: $g(\cdot, \cdot; \theta)$, A sequence of the [PAD] tokens: $\mathbf{e}$ **Output:** Complete (imputed) observation: $\tilde{\mathbf{x}}$ *Phase 1 – Contextual Representation*

```

1: for  $j = 1, 2, \dots, p$  do
2:   if  $\mathbf{m}_j = 1$  then
3:      $\mathbf{t}_j \leftarrow "C_j \text{ is } \mathbf{x}_j"$ 
4:   end if
5:   if  $\mathbf{m}_j = 0$  then
6:      $\mathbf{t}_j \leftarrow "C_j \text{ is } [\text{MASK}]"$   $\triangleright$  replacing missing value with [MASK] token
7:   end if
8: end for
9:  $\mathbf{t} \leftarrow "[\text{CLS}]" \& \mathbf{t}_1 \& \mathbf{t}_2 \& \dots \& \mathbf{t}_p \& "[\text{SEP}]"$   $\triangleright$  Textual encoding
10:  $\mathbf{z} \leftarrow g(\mathbf{t}, \mathbf{e}; \theta) \in \mathbb{R}^D$ 

```

Phase 2 – Context-Driven Imputation

```

11:  $\hat{\mathbf{x}} \leftarrow [0, 0, \dots, 0]$ 
12: for  $j = \{j : \mathbf{m}_j = 0, j = 1, 2, \dots, p\}$  do
13:    $\mathcal{N}_j(\mathbf{x}, k) \subset \{1, 2, \dots, n\}$   $\triangleright$  construct the neighbor index set by Definition 4
14:   if  $j \in I_C$  then
15:      $\hat{\mathbf{x}}_j \leftarrow \frac{1}{k} \sum_{l \in \mathcal{N}_j(\mathbf{x}, k)} \mathbf{x}_j^{(l)}$ 
16:   end if
17:   if  $j \in I_D$  then
18:      $\hat{\mathbf{x}}_j \leftarrow \text{mode}(\mathbf{x}_j^{(l)} : l \in \mathcal{N}_j(\mathbf{x}, k))$ 
19:   end if
20: end for
21:  $\tilde{\mathbf{x}} \leftarrow \mathbf{x} \odot \mathbf{m} + \hat{\mathbf{x}} \odot (1 - \mathbf{m})$ 

```

A.5 EXPERIMENTAL SETTINGS FOR REPRODUCTION

- We run experiments using NVIDIA A10 GPU, and our experimental codes are available with PyTorch and scikit-learn.

Hyper-parameter of DrIM: This demonstrates the generalizability of our proposed model to various tabular datasets. Our implementation codes for the proposed model, DrIM, are provided in the supplementary material. For all tabular datasets, we applied the following hyperparameters uniformly without any additional tuning:

- batch size: 16
- k (the number of neighbors): 5

For fine-tuning the BERT $g(\cdot, \cdot; \theta)$,

- epochs: 5 (with AdamW optimizer (Loshchilov & Hutter, 2017))
- learning rate: 5e-5
- r (the number of re-masking): 3
- τ : 1

A.5.1 DETAILS OF IMPLEMENTING BASELINE MODELS

To compare our proposed method with baseline models, we performed experiments across four missing data mechanisms and five missingness rates. Our reproduced codes for the baseline models

are provided in the supplementary material. Below are the detailed implementations of the baseline methods:

- Mean (Farhangfar et al., 2007): We employ the `SimpleImputer` package ⁷, with the strategy parameter set to ‘mean’.
- kNNI (Troyanskaya et al., 2001): We employ the `KNNImputer` package ⁸ from `scikit-learn` for missing data imputation.
- EM (Nelwamondo et al., 2007): We employ the EM module from `hyperimpute.Plugins` ⁹.
- SoftImpute (Mazumder et al., 2010): We employ the SoftImpute module in `hyperimpute.Plugins` ¹⁰.
- MICE (van Buuren & Groothuis-Oudshoorn, 2011): We employed the `IterativeImputer` package ¹¹ from `scikit-learn`.
- missForest (Stekhoven & Bühlmann, 2012): We employ the missForest module in `hyperimpute.Plugins` ¹².
- Sinkhorn (Muzellec et al., 2020): We employ the Sinkhorn module in `hyperimpute.Plugins` ¹³.
- GAIN (Yoon et al., 2018): We follow the implementations provided in the official repository¹⁴. As the paper does not explicitly discuss the separate handling of categorical and continuous variables, the code treats them simultaneously. Consequently, a rounding process is employed to handle categorical variables afterward.
- VAEAC (Ivanov et al., 2019): We follow the implementations provided in the official repository¹⁵. The authors provided hyperparameters that adequately address both continuous and categorical variables, so we used these without further modification during model fitting.
- MIWAE (Mattei & Frellsen, 2019): The implemented MIWAE code in the official repository¹⁶ was designed for continuous variables only. To accommodate heterogeneous tabular datasets, we treated the conditional distribution of categorical columns as categorical distributions and employed cross-entropy loss for reconstruction.
- not-MIWAE (Ipsen et al., 2021): The implemented not-MIWAE code in the official repository¹⁷ also focused on continuous variables exclusively. To handle categorical variables, we made the same modifications as in MIWAE.
- MIRACLE (Kyono et al., 2021): We utilize the MIRACLE module in `hyperimpute.Plugins` ¹⁸.
- ReMasker (Du et al., 2024a): We follow the implementations provided in official repository ¹⁹.

⁷<https://scikit-learn.org/1.5/modules/generated/sklearn.impute.SimpleImputer.html>

⁸<https://scikit-learn.org/stable/modules/generated/sklearn.impute.KNNImputer.html>

⁹https://github.com/vanderschaarlab/hyperimpute/blob/main/src/hyperimpute/plugins/imputers/plugin_EM.py

¹⁰https://github.com/vanderschaarlab/hyperimpute/blob/main/src/hyperimpute/plugins/imputers/plugin_softimpute.py

¹¹<https://scikit-learn.org/stable/modules/generated/sklearn.impute.IterativeImputer.html>

¹²https://github.com/vanderschaarlab/hyperimpute/blob/main/src/hyperimpute/plugins/imputers/plugin_missforest.py

¹³https://github.com/vanderschaarlab/hyperimpute/blob/main/src/hyperimpute/plugins/imputers/plugin_sinkhorn.py

¹⁴<https://github.com/jsyoon0823/GAIN/tree/master>

¹⁵<https://github.com/tigvarts/vaeac>

¹⁶<https://github.com/pamattei/miwae>

¹⁷<https://github.com/nbip/notMIWAE>

¹⁸https://github.com/vanderschaarlab/hyperimpute/blob/main/src/hyperimpute/plugins/imputers/plugin_miracle.py

¹⁹<https://github.com/tydusky/remasker>

A.6 EVALUATION SETTINGS

Regarding machine learning utility, we conduct classification, model selection, and feature selection tasks for post-imputation evaluation (Please refer to Table 3 for the classification target variable, and Table 4 for detailed machine learning model configuration).

Classification performance (F_1).

1. Train an imputer using a given incomplete training dataset.
2. Impute the incomplete training dataset using trained imputer.
3. Train machine learning models (Logistic Regression (Cox, 1958), Gaussian Naive Bayes, k-nearest neighbors classifier (Altman, 1992), Decision Tree classifier (Wu et al., 2008), and Random Forest classifier (Ho, 1998; Breiman, 2001)) using the imputed dataset.
4. Assess classification prediction performance by averaging the F_1 scores from the complete test dataset from five different classifiers.

Model selection performance (Model).

1. Train an imputer using a given incomplete training dataset.
2. Impute the incomplete training dataset using trained imputer.
3. Train machine learning models (Logistic Regression (Cox, 1958), Gaussian Naive Bayes, k-nearest neighbors classifier (Altman, 1992), Decision Tree classifier (Wu et al., 2008), and Random Forest classifier (Ho, 1998; Breiman, 2001) using both the real complete training dataset and the imputed dataset.
4. Evaluate the classification performance (AUROC) of all trained classifiers on the complete test dataset.
5. Assess model selection performance by comparing the AUROC rank orderings of classifiers trained on the real (original) complete training dataset and those trained on the imputed dataset using Spearman’s Rank Correlation.

Feature selection performance (Feature).

1. Train an imputer using a given incomplete training dataset.
2. Impute the incomplete training dataset using trained imputer.
3. Train a Random Forest classifier (Ho, 1998; Breiman, 2001) using both the real (original) complete training dataset and the imputed dataset.
4. Determine the rank-ordering of important features for both classifiers.
5. Assess feature selection performance by comparing the feature importance rank orderings of classifiers trained on the real complete training dataset and those trained on the imputed dataset using Spearman’s Rank Correlation.

Table 4: **Classifier used to evaluate imputed data quality in machine learning utility.** The names of all parameters used in the description are consistent with those defined in corresponding packages.

Tasks	Model	Description
Classification	Logistic Regression	Package: <code>sklearn.linear_model.LogisticRegression</code> , setting: <code>random_state=0</code> , <code>max_iter=1000</code> , defaulted values
	Gaussian Naive Bayes	Package: <code>sklearn.naive_bayes.GaussianNB</code> , setting: defaulted values
	k-Nearest Neighbors	Package: <code>sklearn.neighbors.KNeighborsClassifier</code> , setting: defaulted values
	Decision Tree	Package: <code>sklearn.tree.DecisionTreeClassifier</code> , setting: <code>random_state=0</code> , defaulted values
	Random Forest	Package: <code>sklearn.ensemble.RandomForestClassifier</code> , setting: <code>random_state=0</code> , defaulted values

A.7 RELATED WORKS

A.7.1 MISSING MECHANISM

Descriptions. Data missingness is a common challenge in research and practical analysis, categorized into three primary missing mechanisms: 1) Missing Completely at Random (MCAR), 2) Missing at Random (MAR), and 3) Missing Not at Random (MNAR).

Under the **MCAR** mechanism, the reason for missingness has no relationship with any data, neither observed nor unobserved. In other words, the likelihood of data being missing is equal across all observations. The primary advantage of MCAR is that it does not introduce bias into the data analysis. However, despite this advantage, data missingness can still reduce the statistical power of the study because of the reduced sample size.

MAR occurs when the probability of missingness is related to the observed data but not the unobserved missing data. Essentially, even though the data is missing, the mechanism assumes that the missingness is explainable by other variables in the dataset. That is, the missingness can be modeled and imputed using the information available in the data, allowing for more accurate analyses despite the missingness.

Lastly, if the missingness is not specified by either MCAR or MAR, it becomes **MNAR**. The MNAR is the most challenging mechanism, as it implies that the missingness is related to the unobserved data itself. In this case, the missing data is systematically different from the observed data, which introduces bias if not properly accounted for. For example, patients with severe symptoms may be less likely to report their health status, making their data missing. MNAR requires sophisticated statistical methods to address, as ignoring or improperly handling it can lead to biased and unreliable results.

Implementations. Following Muzellec et al. (2020); Jarrett et al. (2022); Zhao et al. (2023), we generate the missing value mask for each dataset with three mechanisms in four settings. (MCAR) In the MCAR setting, each value is masked according to the realization of a Bernoulli random variable with a fixed parameter. (MAR) In the MAR setting, for each experiment, a fixed subset of variables that cannot have missing values is sampled. Then, the remaining variables have missing values according to a logistic model with random weights, which takes the non-missing variables as inputs. A bias term is fitted using line search to attain the desired proportion of missing values. (MNAR) Finally, two different mechanisms are implemented in the MNAR setting. The first, MNARL, is identical to the previously described MAR mechanism, but the inputs of the logistic model are then masked by an MCAR mechanism. Hence, the logistic model’s outcome depends on missing values. The second mechanism, MNARQ, samples a subset of variables whose values in the lower and upper p th percentiles are masked according to a Bernoulli random variable, and the values in-between are left not missing.

A.8 ADDITIONAL EXPERIMENTS

We conduct extensive additional experiments using a variety of pre-trained language models, including both Autoencoding (BERT, RoBERTa) and Autoregressive (GPT-2, GPT-NEO, LLaMA) models, across different parameter scales.

A.8.1 CONFIGURATION OF LANGUAGE MODELS

Table 5: The language model used to generate contextualized representations. ‘#Params’ represents the number of parameters as defined in the respective language model specifications.

Type	model	#Params
Autoencoding	BERT _{BASE} (Devlin et al., 2019)	0.11B
	RoBERTa (Liu et al., 2019)	0.12B
	BERT _{LARGE} (Devlin et al., 2019)	0.34B
Autoregressive	GPT-2 (Radford et al., 2019)	1.56B
	GPT-NEO (Black et al., 2021)	1.32B
	LLaMA (Touvron et al., 2023)	6.61B

A.8.2 EXPERIMENT RESULTS

Table 6: **Machine learning utility** at 0.3 missingness rate. Results using a pretrained language model without fine-tuning. $\uparrow$ denotes that higher values are better. The means and standard errors of the mean across 10 datasets and 10 repeated experiments are reported.

Type	model	MCAR			MAR		
		$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Autoencoding	BERT _{BASE}	.638 $\pm$.024	.517 $\pm$.037	.920 $\pm$.006	.640 $\pm$.024	.617 $\pm$.041	.894 $\pm$.010
	RoBERTa	.652 $\pm$.024	.538 $\pm$.038	.922 $\pm$.006	.652 $\pm$.024	.628 $\pm$.037	.907 $\pm$.007
	BERT _{LARGE}	.647 $\pm$.025	.625 $\pm$.031	.898 $\pm$.008	.652 $\pm$.025	.641 $\pm$.030	.909 $\pm$.007
Autoregressive	GPT-2	.658 $\pm$.024	.620 $\pm$.032	.908 $\pm$.008	.665 $\pm$.025	.685 $\pm$.032	.907 $\pm$.007
	GPT-NEO	.659 $\pm$.025	.606 $\pm$.032	.922 $\pm$.006	.665 $\pm$.025	.688 $\pm$.033	.916 $\pm$.006
	LLaMA	.657 $\pm$.024	.659 $\pm$.030	.911 $\pm$.009	.664 $\pm$.024	.737 $\pm$.026	.890 $\pm$.010

Type	model	MNARL			MNARQ		
		$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Autoencoding	BERT _{BASE}	.625 $\pm$.024	.496 $\pm$.043	.886 $\pm$.009	.658 $\pm$.025	.694 $\pm$.029	.913 $\pm$.007
	RoBERTa	.641 $\pm$.023	.561 $\pm$.042	.901 $\pm$.008	.672 $\pm$.024	.673 $\pm$.028	.918 $\pm$.006
	BERT _{LARGE}	.635 $\pm$.025	.596 $\pm$.033	.886 $\pm$.009	.670 $\pm$.025	.713 $\pm$.025	.911 $\pm$.007
Autoregressive	GPT-2	.655 $\pm$.024	.607 $\pm$.036	.906 $\pm$.007	.677 $\pm$.025	.741 $\pm$.024	.921 $\pm$.006
	GPT-NEO	.653 $\pm$.024	.640 $\pm$.032	.908 $\pm$.006	.675 $\pm$.025	.710 $\pm$.026	.920 $\pm$.005
	LLaMA	.653 $\pm$.024	.666 $\pm$.032	.900 $\pm$.008	.676 $\pm$.024	.726 $\pm$.025	.886 $\pm$.012

A.8.3 TIME EFFICIENCY

Table 7: **Time Efficiency**. The reported unit is second(s). Results using a pretrained language model without fine-tuning. The means and standard errors of the mean across 10 datasets and 10 repeated experiments are reported.

Time (s)	model	MCAR	MAR	MNARL	MNARQ
Autoencoding	BERT _{BASE}	93.240 $\pm$ 9.887	86.160 $\pm$ 9.089	94.030 $\pm$ 9.767	97.860 $\pm$ 10.399
	RoBERTa	92.780 $\pm$ 9.814	84.190 $\pm$ 8.809	96.020 $\pm$ 10.274	94.770 $\pm$ 10.136
	BERT _{LARGE}	122.380 $\pm$ 13.025	113.220 $\pm$ 12.186	127.010 $\pm$ 13.699	128.810 $\pm$ 14.126
Autoregressive	GPT-2	292.780 $\pm$ 32.322	272.340 $\pm$ 30.304	298.500 $\pm$ 33.329	284.930 $\pm$ 31.403
	GPT-NEO	237.670 $\pm$ 25.515	216.765 $\pm$ 24.932	233.190 $\pm$ 24.546	235.620 $\pm$ 25.319
	LLaMA	1196.240 $\pm$ 132.644	1276.640 $\pm$ 144.409	1196.470 $\pm$ 132.269	1277.320 $\pm$ 143.210

A.9 DETAILED EXPERIMENTAL RESULT

A.9.1 IMPUTATION FIDELITY

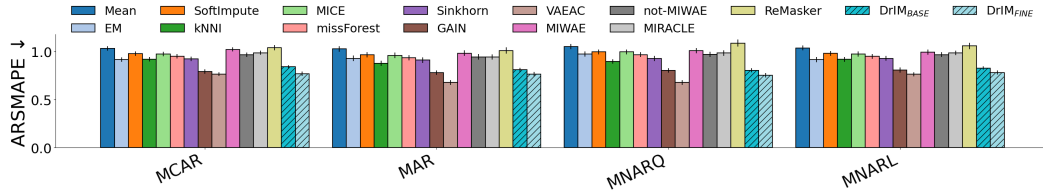


Figure 4: Imputation fidelity at **0.2** missingness rate. The missing mechanism corresponding to the figure is indicated below the figure. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported.

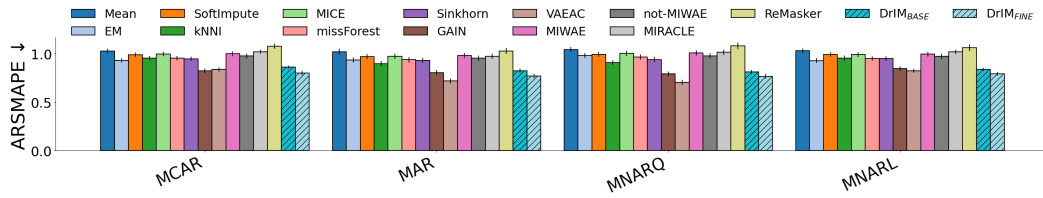


Figure 5: Imputation fidelity at **0.4** missingness rate. The missing mechanism corresponding to the figure is indicated below the figure. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported.

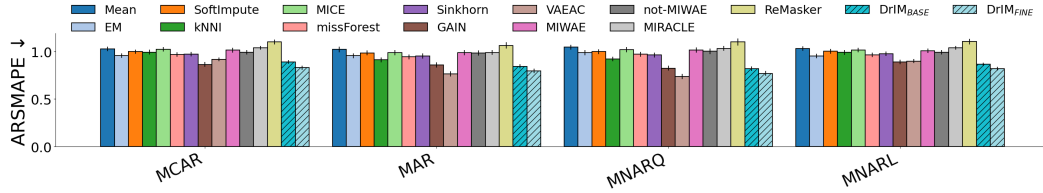


Figure 6: Imputation fidelity at **0.6** missingness rate. The missing mechanism corresponding to the figure is indicated below the figure. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported.

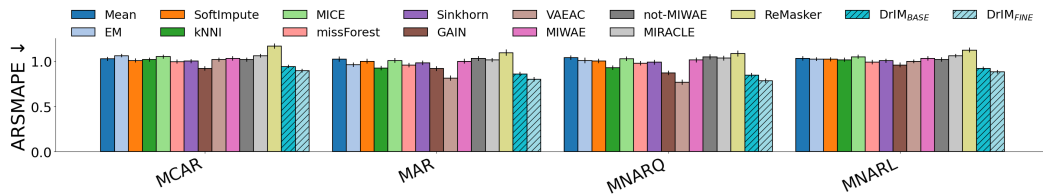


Figure 7: Imputation fidelity at **0.8** missingness rate. The missing mechanism corresponding to the figure is indicated below the figure. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported.

A.9.2 MACHINE LEARNING UTILITY

Table 8: Machine learning utility at **0.2** missingness rate. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. The best value is bolded, and the second best is underlined.

model	MCAR			MAR			MNARL			MNARQ		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.646 $\pm$.026	.543 $\pm$.037	.936 $\pm$.005	.642 $\pm$.026	.710 $\pm$.030	.863 $\pm$.013	.629 $\pm$.026	.623 $\pm$.034	.904 $\pm$.009	.668 $\pm$.026	.651 $\pm$.034	.899 $\pm$.009
EM	.657 $\pm$.025	.523 $\pm$.032	.934 $\pm$.007	.641 $\pm$.026	.707 $\pm$.037	.912 $\pm$.010	.640 $\pm$.026	.619 $\pm$.038	.886 $\pm$.012	.669 $\pm$.025	.623 $\pm$.032	.916 $\pm$.008
SoftImpute	.650 $\pm$.026	.530 $\pm$.035	.939 $\pm$.005	.643 $\pm$.026	.694 $\pm$.032	.905 $\pm$.009	.631 $\pm$.026	.576 $\pm$.036	.903 $\pm$.009	.669 $\pm$.026	.670 $\pm$.031	.932 $\pm$.006
kNNI	.654 $\pm$.025	.572 $\pm$.032	.940 $\pm$.008	.649 $\pm$.025	.705 $\pm$.032	.923 $\pm$.010	.636 $\pm$.026	.605 $\pm$.039	.899 $\pm$.012	.672 $\pm$.026	.646 $\pm$.032	.930 $\pm$.007
MICE	.656 $\pm$.026	.489 $\pm$.035	.946 $\pm$.005	.648 $\pm$.026	.690 $\pm$.034	.912 $\pm$.009	.637 $\pm$.026	.601 $\pm$.036	.904 $\pm$.010	.673 $\pm$.026	.633 $\pm$.033	.933 $\pm$.007
missForest	.651 $\pm$.026	.557 $\pm$.033	.925 $\pm$.007	.646 $\pm$.026	.705 $\pm$.032	.903 $\pm$.009	.634 $\pm$.026	.595 $\pm$.039	.894 $\pm$.011	.671 $\pm$.026	.659 $\pm$.033	.920 $\pm$.007
Sinkhorn	.651 $\pm$.026	.563 $\pm$.035	.932 $\pm$.008	.649 $\pm$.025	.696 $\pm$.033	.911 $\pm$.009	.635 $\pm$.026	.617 $\pm$.037	.900 $\pm$.011	.671 $\pm$.026	.649 $\pm$.033	.918 $\pm$.008
GAIN	.669 $\pm$.025	.659 $\pm$.022	.935 $\pm$.005	.673 $\pm$.025	.682 $\pm$.025	.931 $\pm$.006	.662 $\pm$.024	.630 $\pm$.032	.914 $\pm$.008	.685 $\pm$.025	.682 $\pm$.026	.931 $\pm$.005
VAEAC	.611 $\pm$.025	.634 $\pm$.019	.857 $\pm$.004	.610 $\pm$.025	.651 $\pm$.017	.853 $\pm$.005	.609 $\pm$.025	.626 $\pm$.019	.851 $\pm$.005	.614 $\pm$.024	.679 $\pm$.014	.847 $\pm$.008
MIWAE	.632 $\pm$.026	.581 $\pm$.034	.931 $\pm$.007	.641 $\pm$.026	.669 $\pm$.034	.902 $\pm$.009	.629 $\pm$.026	.578 $\pm$.039	.896 $\pm$.011	.667 $\pm$.026	.645 $\pm$.034	.929 $\pm$.006
not-MIWAE	.647 $\pm$.025	.593 $\pm$.032	.939 $\pm$.006	.643 $\pm$.026	.663 $\pm$.035	.904 $\pm$.010	.628 $\pm$.026	.597 $\pm$.037	.898 $\pm$.010	.669 $\pm$.026	.669 $\pm$.032	.932 $\pm$.006
MIRACLE	.646 $\pm$.025	.474 $\pm$.038	.937 $\pm$.006	.646 $\pm$.026	.609 $\pm$.039	.911 $\pm$.009	.630 $\pm$.026	.549 $\pm$.041	.889 $\pm$.012	.670 $\pm$.026	.629 $\pm$.033	.923 $\pm$.008
ReMasker	.648 $\pm$.026	.493 $\pm$.037	.907 $\pm$.010	.643 $\pm$.026	.640 $\pm$.037	.888 $\pm$.014	.630 $\pm$.027	.523 $\pm$.041	.870 $\pm$.013	.666 $\pm$.026	.599 $\pm$.036	.919 $\pm$.008
DrIM _{BASE}	.665 $\pm$.025	.692 $\pm$.026	.949 $\pm$.004	<u>.673</u> $\pm$.024	<u>.732</u> $\pm$.029	<u>.934</u> $\pm$.006	.660 $\pm$.024	.658 $\pm$.031	.925 $\pm$.007	<u>.682</u> $\pm$.025	<u>.771</u> $\pm$.026	<u>.945</u> $\pm$.005
DrIM _{FINE}	.678 $\pm$.025	<u>.690</u> $\pm$.029	<u>.948</u> $\pm$.004	.684 $\pm$.025	.773 $\pm$.023	.948 $\pm$.004	.673 $\pm$.025	.711 $\pm$.026	.938 $\pm$.006	.689 $\pm$.025	.784 $\pm$.021	.948 $\pm$.004

Table 9: Machine learning utility at **0.3** missingness rate. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. The best value is bolded, and the second best is underlined.

model	MCAR			MAR			MNARL			MNARQ		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.581 $\pm$.027	.289 $\pm$.046	.839 $\pm$.017	.591 $\pm$.027	.548 $\pm$.041	.744 $\pm$.021	.560 $\pm$.027	.428 $\pm$.041	.797 $\pm$.017	.620 $\pm$.028	.417 $\pm$.052	.771 $\pm$.021
kNNI	.595 $\pm$.026	.292 $\pm$.047	.895 $\pm$.014	.601 $\pm$.027	.560 $\pm$.043	.868 $\pm$.014	.575 $\pm$.027	.419 $\pm$.049	.843 $\pm$.016	.629 $\pm$.027	.448 $\pm$.048	.877 $\pm$.013
EM	.599 $\pm$.027	.244 $\pm$.048	.873 $\pm$.015	.598 $\pm$.027	.531 $\pm$.043	.833 $\pm$.016	.580 $\pm$.027	.407 $\pm$.048	.804 $\pm$.019	.626 $\pm$.028	.412 $\pm$.050	.835 $\pm$.018
SoftImpute	.583 $\pm$.027	.233 $\pm$.044	.866 $\pm$.013	.590 $\pm$.027	.554 $\pm$.039	.845 $\pm$.014	.562 $\pm$.027	.342 $\pm$.045	.818 $\pm$.015	.622 $\pm$.028	.403 $\pm$.050	.871 $\pm$.013
MICE	.593 $\pm$.027	.247 $\pm$.046	.892 $\pm$.013	.598 $\pm$.027	.543 $\pm$.043	.859 $\pm$.014	.571 $\pm$.027	.339 $\pm$.050	.825 $\pm$.017	.628 $\pm$.028	.410 $\pm$.051	.876 $\pm$.014
missForest	.588 $\pm$.027	.229 $\pm$.046	.858 $\pm$.015	.596 $\pm$.027	.548 $\pm$.044	.833 $\pm$.015	.569 $\pm$.027	.366 $\pm$.048	.809 $\pm$.016	.625 $\pm$.028	.401 $\pm$.048	.857 $\pm$.014
Sinkhorn	.589 $\pm$.027	.242 $\pm$.043	.868 $\pm$.015	.599 $\pm$.027	.545 $\pm$.044	.845 $\pm$.015	.569 $\pm$.027	.377 $\pm$.049	.821 $\pm$.017	.626 $\pm$.028	.448 $\pm$.045	.840 $\pm$.019
GAIN	.629 $\pm$.025	.473 $\pm$.034	.856 $\pm$.012	<u>.640</u> $\pm$.025	.602 $\pm$.029	.877 $\pm$.011	.621 $\pm$.025	.446 $\pm$.040	.843 $\pm$.016	<u>.659</u> $\pm$.026	.567 $\pm$.037	.875 $\pm$.013
VAEAC	.604 $\pm$.025	<u>.558</u> $\pm$.022	.832 $\pm$.007	.606 $\pm$.025	.612 $\pm$.018	.827 $\pm$.007	.598 $\pm$.025	.553 $\pm$.024	.816 $\pm$.012	.610 $\pm$.025	.593 $\pm$.020	.823 $\pm$.012
MIWAE	.577 $\pm$.027	.245 $\pm$.052	.864 $\pm$.014	.588 $\pm$.027	.519 $\pm$.048	.834 $\pm$.014	.556 $\pm$.028	.369 $\pm$.053	.821 $\pm$.016	.618 $\pm$.028	.405 $\pm$.054	.863 $\pm$.014
not-MIWAE	.580 $\pm$.027	.234 $\pm$.049	.869 $\pm$.015	.587 $\pm$.027	.540 $\pm$.043	.839 $\pm$.013	.558 $\pm$.027	.394 $\pm$.049	.818 $\pm$.017	.622 $\pm$.028	.438 $\pm$.049	.874 $\pm$.013
MIRACLE	.576 $\pm$.027	.167 $\pm$.048	.876 $\pm$.013	.597 $\pm$.027	.473 $\pm$.046	.851 $\pm$.014	.564 $\pm$.027	.264 $\pm$.051	.825 $\pm$.016	.625 $\pm$.028	.333 $\pm$.053	.865 $\pm$.016
ReMasker	.583 $\pm$.029	.229 $\pm$.046	.823 $\pm$.017	.592 $\pm$.027	.480 $\pm$.046	.817 $\pm$.019	<u>.566</u> $\pm$.028	.279 $\pm$.050	.782 $\pm$.021	.617 $\pm$.028	.340 $\pm$.050	.835 $\pm$.016
DrIM _{BASE}	<u>.638</u> $\pm$.024	.517 $\pm$.037	<u>.920</u> $\pm$.006	<u>.640</u> $\pm$.024	<u>.617</u> $\pm$.041	.894 $\pm$.010	<u>.625</u> $\pm$.024	.496 $\pm$.043	.886 $\pm$.009	.658 $\pm$.025	.694 $\pm$.029	.913 $\pm$.007
DrIM _{FINE}	.656 $\pm$.025	.616 $\pm$.032	.921 $\pm$.006	.659 $\pm$.025	.658 $\pm$.032	.905 $\pm$.007	.653 $\pm$.025	.553 $\pm$.040	.915 $\pm$.006	.673 $\pm$.025	<u>.691</u> $\pm$.030	<u>.910</u> $\pm$.007

Table 10: Machine learning utility at **0.4** missingness rate. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. The best value is bolded, and the second best is underlined.

model	MCAR			MAR			MNARL			MNARQ		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.496 $\pm$.030	.062 $\pm$.046	.672 $\pm$.037	.533 $\pm$.028	.370 $\pm$.047	.607 $\pm$.029	.493 $\pm$.029	.201 $\pm$.050	.662 $\pm$.031	.561 $\pm$.030	.293 $\pm$.057	.632 $\pm$.030
EM	.532 $\pm$.028	.093 $\pm$.052	.819 $\pm$.018	.540 $\pm$.028	.369 $\pm$.050	.780 $\pm$.020	.522 $\pm$.028	.179 $\pm$.050	.733 $\pm$.024	.578 $\pm$.030	.302 $\pm$.056	.767 $\pm$.024
SoftImpute	.491 $\pm$.030	-.080 $\pm$.045	.740 $\pm$.025	.534 $\pm$.029	.285 $\pm$.049	.743 $\pm$.024	.491 $\pm$.029	.118 $\pm$.052	.719 $\pm$.022	.563 $\pm$.030	.240 $\pm$.058	.787 $\pm$.023
kNNI	.529 $\pm$.028	.087 $\pm$.053	.849 $\pm$.016	.547 $\pm$.028	.343 $\pm$.051	.806 $\pm$.018	.521 $\pm$.027	.238 $\pm$.055	.779 $\pm$.021	.576 $\pm$.030	.271 $\pm$.055	.806 $\pm$.021
MICE	.509 $\pm$.030	-.032 $\pm$.049	.851 $\pm$.015	.539 $\pm$.029	.319 $\pm$.053	.817 $\pm$.017	.501 $\pm$.029	.084 $\pm$.052	.774 $\pm$.020	.569 $\pm$.031	.263 $\pm$.057	.816 $\pm$.020
missForest	.509 $\pm$.029	-.001 $\pm$.044	.791 $\pm$.021	.541 $\pm$.028	.329 $\pm$.050	.775 $\pm$.019	.503 $\pm$.029	.182 $\pm$.049	.731 $\pm$.022	.570 $\pm$.030	.266 $\pm$.056	.779 $\pm$.022
Sinkhorn	.511 $\pm$.029	.001 $\pm$.047	.798 $\pm$.021	.540 $\pm$.028	.341 $\pm$.048	.776 $\pm$.021	.505 $\pm$.028	.173 $\pm$.049	.748 $\pm$.022	.572 $\pm$.030	.271 $\pm$.056	.748 $\pm$.027
GAIN	.587 $\pm$.026	.320 $\pm$.044	.725 $\pm$.028	<u>.612</u> $\pm$.026	.475 $\pm$.036	.781 $\pm$.019	.583 $\pm$.026	.250 $\pm$.043	.726 $\pm$.022	<u>.638</u> $\pm$.026	.461 $\pm$.040	.794 $\pm$.019
VAEAC	.599 $\pm$.025	<u>.488</u>	.789 $\pm$.013	.598 $\pm$.025	<u>.521</u> $\pm$.025	.790 $\pm$.012	<u>.592</u> $\pm$.025	.495 $\pm$.025	.788 $\pm$.014	.603 $\pm$.025	<u>.585</u> $\pm$.023	.782 $\pm$.016
MIWAE	.491 $\pm$.030	-.041 $\pm$.047	.791 $\pm$.021	.530 $\pm$.028	.317 $\pm$.051	.781 $\pm$.019	.487 $\pm$.029	.083 $\pm$.049	.739 $\pm$.021	.560 $\pm$.030	.236 $\pm$.056	.790 $\pm$.021
not-MIWAE	.490 $\pm$.030	-.074 $\pm$.050	.798 $\pm$.020	.529 $\pm$.028	.347 $\pm$.049	.770 $\pm$.020	.486 $\pm$.029	.103 $\pm$.054	.746 $\pm$.019	.559 $\pm$.031	.246 $\pm$.058	.780 $\pm$.022
MIRACLE	.499 $\pm$.029	-.055 $\pm$.049	.812 $\pm$.015	.535 $\pm$.028	.277 $\pm$.058	.769 $\pm$.021	.493 $\pm$.028	.029 $\pm$.051	.737 $\pm$.024	.566 $\pm$.030	.203 $\pm$.054	.797 $\pm$.023
ReMasker	.490 $\pm$.031	.081 $\pm$.050	.750 $\pm$.026	.535 $\pm$.029	.341 $\pm$.048	.718 $\pm$.028	.499 $\pm$.029	.104 $\pm$.054	.703 $\pm$.027	.557 $\pm$.031	.249 $\pm$.050	.759 $\pm$.025
DrIMBASE	<u>.602</u> $\pm$.024	.301 $\pm$.048	<u>.853</u> $\pm$.012	.607 $\pm$.024	.452 $\pm$.044	<u>.837</u> $\pm$.015	.591 $\pm$.023	.326 $\pm$.048	<u>.832</u> $\pm$.014	.629 $\pm$.025	.534 $\pm$.037	<u>.850</u> $\pm$.012
DrIMFINE	.632 $\pm$.025	.503 $\pm$.038	.889 $\pm$.011	.639 $\pm$.025	.556 $\pm$.035	.853 $\pm$.012	.625 $\pm$.024	<u>.481</u> $\pm$.042	.864 $\pm$.010	.655 $\pm$.025	.618 $\pm$.033	.853 $\pm$.012

Table 11: Machine learning utility at **0.6** missingness rate. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. The best value is bolded, and the second best is underlined.

model	MCAR			MAR			MNARL			MNARQ		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.377 $\pm$.031	-.186 $\pm$.051	.422 $\pm$.047	.433 $\pm$.030	.012 $\pm$.059	.411 $\pm$.037	.380 $\pm$.030	-.138 $\pm$.048	.427 $\pm$.043	.464 $\pm$.034	.034 $\pm$.063	.450 $\pm$.041
EM	.429 $\pm$.030	-.255 $\pm$.047	.669 $\pm$.025	.461 $\pm$.029	.024 $\pm$.054	.626 $\pm$.029	.432 $\pm$.029	-.175 $\pm$.049	.569 $\pm$.033	.494 $\pm$.033	.024 $\pm$.059	.567 $\pm$.034
SoftImpute	.363 $\pm$.030	-.277 $\pm$.053	.445 $\pm$.045	.432 $\pm$.031	-.027 $\pm$.060	.469 $\pm$.043	.374 $\pm$.030	-.226 $\pm$.048	.472 $\pm$.041	.462 $\pm$.034	-.006 $\pm$.064	.567 $\pm$.037
kNNI	.441 $\pm$.029	-.151 $\pm$.058	.749 $\pm$.020	.458 $\pm$.030	.033 $\pm$.057	.677 $\pm$.024	.439 $\pm$.029	-.100 $\pm$.054	.609 $\pm$.027	.486 $\pm$.033	.074 $\pm$.061	.659 $\pm$.031
MICE	.361 $\pm$.031	-.295 $\pm$.052	.732 $\pm$.019	.441 $\pm$.031	-.001 $\pm$.055	.684 $\pm$.026	.378 $\pm$.031	-.247 $\pm$.052	.624 $\pm$.028	.468 $\pm$.034	.029 $\pm$.061	.673 $\pm$.031
missForest	.379 $\pm$.031	-.257 $\pm$.049	.626 $\pm$.029	.447 $\pm$.031	-.038 $\pm$.057	.610 $\pm$.031	.389 $\pm$.030	-.229 $\pm$.050	.546 $\pm$.030	.475 $\pm$.034	.005 $\pm$.064	.609 $\pm$.034
Sinkhorn	.389 $\pm$.030	-.230 $\pm$.052	.542 $\pm$.037	.448 $\pm$.030	-.025 $\pm$.058	.572 $\pm$.033	.399 $\pm$.030	-.214 $\pm$.052	.539 $\pm$.031	.479 $\pm$.033	.034 $\pm$.064	.560 $\pm$.039
GAIN	.499 $\pm$.027	-.073 $\pm$.053	.482 $\pm$.040	.525 $\pm$.027	.172 $\pm$.049	.585 $\pm$.030	.494 $\pm$.025	-.058 $\pm$.049	.456 $\pm$.035	.575 $\pm$.027	.237 $\pm$.054	.593 $\pm$.035
VAEAC	.565 $\pm$.024	<u>.391</u> $\pm$.028	.693 $\pm$.021	<u>.574</u> $\pm$.025	.421 $\pm$.033	.723 $\pm$.019	<u>.549</u> $\pm$.025	.339 $\pm$.033	.643 $\pm$.026	<u>.587</u> $\pm$.025	<u>.462</u> $\pm$.032	<u>.693</u> $\pm$.026
MIWAE	.357 $\pm$.031	-.313 $\pm$.048	.669 $\pm$.028	.433 $\pm$.031	-.028 $\pm$.056	.638 $\pm$.028	.376 $\pm$.031	-.299 $\pm$.047	.583 $\pm$.030	.460 $\pm$.034	.032 $\pm$.064	.656 $\pm$.033
not-MIWAE	.355 $\pm$.031	-.290 $\pm$.049	.625 $\pm$.032	.420 $\pm$.031	-.056 $\pm$.057	.583 $\pm$.031	.367 $\pm$.030	-.298 $\pm$.049	.556 $\pm$.032	.456 $\pm$.034	-.020 $\pm$.063	.617 $\pm$.033
MIRACLE	.394 $\pm$.029	-.253 $\pm$.049	.634 $\pm$.022	.448 $\pm$.030	-.089 $\pm$.058	.567 $\pm$.033	.393 $\pm$.028	-.217 $\pm$.051	.521 $\pm$.031	.480 $\pm$.033	.015 $\pm$.059	.612 $\pm$.032
ReMasker	.378 $\pm$.032	-.221 $\pm$.052	.585 $\pm$.032	.441 $\pm$.031	-.018 $\pm$.056	.521 $\pm$.040	.386 $\pm$.030	-.253 $\pm$.051	.509 $\pm$.041	.469 $\pm$.034	.048 $\pm$.061	.586 $\pm$.036
DrIMBASE	.539 $\pm$.024	.104 $\pm$.051	.708 $\pm$.024	.545 $\pm$.024	<u>.231</u> $\pm$.052	.684 $\pm$.024	.522 $\pm$.022	.102 $\pm$.052	<u>.652</u> $\pm$.029	.583 $\pm$.026	.266 $\pm$.050	.683 $\pm$.028
DrIMFINE	.588 $\pm$.025	.349 $\pm$.044	<u>.743</u> $\pm$.019	.589 $\pm$.024	.373 $\pm$.040	.734 $\pm$.026	.573 $\pm$.024	<u>.313</u> $\pm$.045	.730 $\pm$.021	.614 $\pm$.026	<u>.434</u> $\pm$.042	.748 $\pm$.021

Table 12: Machine learning utility at **0.8** missingness rate. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. The best value is bolded, and the second best is underlined.

model	MCAR			MAR			MNARL			MNARQ		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.322 $\pm$.030	-.180 $\pm$.049	.247 $\pm$.052	.353 $\pm$.031	-.017 $\pm$.056	.278 $\pm$.039	.311 $\pm$.030	-.064 $\pm$.051	.294 $\pm$.043	.413 $\pm$.034	.043 $\pm$.061	.346 $\pm$.041
EM	.366 $\pm$.029	-.186 $\pm$.053	.381 $\pm$.042	.386 $\pm$.029	-.101 $\pm$.060	.418 $\pm$.039	.362 $\pm$.028	-.200 $\pm$.050	.345 $\pm$.039	.441 $\pm$.033	-.039 $\pm$.061	.402 $\pm$.038
SoftImpute	.315 $\pm$.030	-.311 $\pm$.048	.263 $\pm$.048	.362 $\pm$.031	-.077 $\pm$.057	.322 $\pm$.040	.311 $\pm$.030	-.165 $\pm$.050	.312 $\pm$.044	.415 $\pm$.035	.050 $\pm$.060	.376 $\pm$.041
kNNI	.396 $\pm$.031	-.199 $\pm$.055	.585 $\pm$.030	.386 $\pm$.032	-.120 $\pm$.053	.549 $\pm$.031	.386 $\pm$.031	-.061 $\pm$.054	<u>.440</u> $\pm$.036	.431 $\pm$.035	-.026 $\pm$.060	.507 $\pm$.039
MICE	.292 $\pm$.031	-.331 $\pm$.049	.476 $\pm$.030	.369 $\pm$.033	-.109 $\pm$.061	.502 $\pm$.035	.295 $\pm$.031	-.260 $\pm$.049	.381 $\pm$.040	.417 $\pm$.036	-.040 $\pm$.060	.494 $\pm$.039
missForest	.315 $\pm$.031	-.306 $\pm$.047	.331 $\pm$.038	.375 $\pm$.032	-.138 $\pm$.055	.426 $\pm$.039	.314 $\pm$.031	-.215 $\pm$.051	.314 $\pm$.042	.422 $\pm$.035	-.030 $\pm$.059	.422 $\pm$.041
Sinkhorn	.330 $\pm$.030	-.211 $\pm$.054	.272 $\pm$.040	.373 $\pm$.031	-.010 $\pm$.055	.389 $\pm$.039	.324 $\pm$.029	-.099 $\pm$.055	.299 $\pm$.039	.425 $\pm$.034	-.063 $\pm$.058	.401 $\pm$.042
GAIN	.411 $\pm$.026	-.161 $\pm$.054	.298 $\pm$.042	.438 $\pm$.027	-.001 $\pm$.054	.364 $\pm$.040	.391 $\pm$.026	-.136 $\pm$.052	.244 $\pm$.040	.514 $\pm$.029	.155 $\pm$.059	.451 $\pm$.041
VAEAC	.505 $\pm$.024	.185 $\pm$.039	<u>.534</u> $\pm$.030	.536 $\pm$.025	.316 $\pm$.036	<u>.611</u> $\pm$.026	<u>.492</u> $\pm$.024	<u>.217</u> $\pm$.041	.488 $\pm$.033	.558 $\pm$.026	.322 $\pm$.037	<u>.660</u> $\pm$.032
MIWAE	.287 $\pm$.031	-.322 $\pm$.049	.397 $\pm$.039	.364 $\pm$.032	-.116 $\pm$.059	.454 $\pm$.036	.299 $\pm$.031	-.283 $\pm$.046	.348 $\pm$.038	.409 $\pm$.035	-.032 $\pm$.060	.487 $\pm$.038
not-MIWAE	.288 $\pm$.030	-.364 $\pm$.042	.357 $\pm$.035	.347 $\pm$.032	-.146 $\pm$.056	.406 $\pm$.039	.288 $\pm$.030	-.332 $\pm$.044	.335 $\pm$.037	.405 $\pm$.035	-.024 $\pm$.057	.473 $\pm$.033
MIRACLE	.320 $\pm$.028	-.273 $\pm$.049	.368 $\pm$.032	.369 $\pm$.030	-.183 $\pm$.058	.367 $\pm$.040	.316 $\pm$.028	-.216 $\pm$.050	.290 $\pm$.036	.419 $\pm$.033	-.097 $\pm$.054	.420 $\pm$.039
ReMasker	.304 $\pm$.028	-.283 $\pm$.053	.299 $\pm$.049	.374 $\pm$.031	-.140 $\pm$.055	.364 $\pm$.042	.306 $\pm$.029	-.285 $\pm$.051	.286 $\pm$.045	.418 $\pm$.035	-.034 $\pm$.058	.418 $\pm$.037
DrIMBASE	.471 $\pm$.023	.027 $\pm$.050	.272 $\pm$.043	.477 $\pm$.023	.072 $\pm$.054	.446 $\pm$.041	<u>.449</u> $\pm$.022	.049 $\pm$.051	.286 $\pm$.043	.515 $\pm$.026	.151 $\pm$.060	.545 $\pm$.033
DrIMFINE	<u>.503</u> $\pm$.025	<u>.176</u> $\pm$.055	.309 $\pm$.039	<u>.534</u> $\pm$.024	<u>.257</u> $\pm$.047	<u>.604</u> $\pm$.035	.492 $\pm$.024	.235 $\pm$.054	.287 $\pm$.042	<u>.572</u> $\pm$.025	<u>.281</u> $\pm$.052	.616 $\pm$.031

A.9.3 MACHINE LEARNING UTILITY FOR EACH DATASET

Table 13: Machine learning utility for each dataset under **MCAR** at 0.3 missingness. The means and the standard errors of the mean across 10 repeated experiments are reported. $\uparrow$ denotes higher is better.

model	abalone			anuran		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.125 $\pm$.002	-.587 $\pm$.059	.864 $\pm$.015	.858 $\pm$.002	.172 $\pm$.113	.836 $\pm$.012
kNNI	.135 $\pm$.003	-.475 $\pm$.069	.983 $\pm$.005	.832 $\pm$.003	.105 $\pm$.106	.496 $\pm$.008
EM	.128 $\pm$.003	-.505 $\pm$.074	.962 $\pm$.012	.841 $\pm$.003	-.133 $\pm$.091	.457 $\pm$.011
SoftImpute	.118 $\pm$.004	-.523 $\pm$.054	.967 $\pm$.010	.855 $\pm$.003	.125 $\pm$.126	.626 $\pm$.026
MICE	.120 $\pm$.003	-.542 $\pm$.048	.974 $\pm$.007	.858 $\pm$.004	.105 $\pm$.111	.545 $\pm$.014
missForest	.121 $\pm$.003	-.507 $\pm$.045	.976 $\pm$.006	.850 $\pm$.003	.080 $\pm$.109	.485 $\pm$.012
Sinkhorn	.125 $\pm$.003	-.547 $\pm$.069	.967 $\pm$.010	.843 $\pm$.003	.015 $\pm$.094	.473 $\pm$.009
GAIN	.196 $\pm$.003	.019 $\pm$.150	.850 $\pm$.014	.902 $\pm$.011	.452 $\pm$.059	.772 $\pm$.032
VAEAC	.132 $\pm$.002	.322 $\pm$.087	.821 $\pm$.013	.892 $\pm$.000	.535 $\pm$.070	.872 $\pm$.004
MIWAE	.113 $\pm$.003	-.578 $\pm$.047	.967 $\pm$.016	.855 $\pm$.003	.225 $\pm$.099	.517 $\pm$.019
not-MIWAE	.114 $\pm$.003	-.607 $\pm$.049	.957 $\pm$.011	.854 $\pm$.003	.201 $\pm$.121	.521 $\pm$.014
MIRACLE	.126 $\pm$.003	-.446 $\pm$.051	.967 $\pm$.009	.821 $\pm$.005	-.188 $\pm$.093	.538 $\pm$.022
ReMasker	.119 $\pm$.004	-.543 $\pm$.051	.886 $\pm$.022	.837 $\pm$.003	.045 $\pm$.090	.479 $\pm$.010
DrIM _{BASE}	.200 $\pm$.002	.571 $\pm$.142	.893 $\pm$.021	.926 $\pm$.001	.262 $\pm$.092	.947 $\pm$.006
DrIM _{FINE}	.203 $\pm$.003	.534 $\pm$.109	.926 $\pm$.012	.970 $\pm$.007	.384 $\pm$.112	.938 $\pm$.005

model	banknote			breast		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.863 $\pm$.006	.731 $\pm$.056	1.000 $\pm$.000	.896 $\pm$.006	.198 $\pm$.119	.872 $\pm$.011
kNNI	.869 $\pm$.006	.641 $\pm$.095	1.000 $\pm$.000	.883 $\pm$.010	-.007 $\pm$.107	.878 $\pm$.014
EM	.912 $\pm$.005	.525 $\pm$.075	1.000 $\pm$.000	.885 $\pm$.008	.004 $\pm$.074	.876 $\pm$.015
SoftImpute	.886 $\pm$.006	.438 $\pm$.055	1.000 $\pm$.000	.881 $\pm$.008	.147 $\pm$.082	.865 $\pm$.015
MICE	.906 $\pm$.006	.493 $\pm$.056	1.000 $\pm$.000	.888 $\pm$.007	.023 $\pm$.089	.873 $\pm$.016
missForest	.895 $\pm$.005	.498 $\pm$.047	1.000 $\pm$.000	.880 $\pm$.007	-.072 $\pm$.117	.892 $\pm$.009
Sinkhorn	.882 $\pm$.004	.483 $\pm$.042	1.000 $\pm$.000	.886 $\pm$.008	.035 $\pm$.111	.885 $\pm$.014
GAIN	.909 $\pm$.012	.577 $\pm$.072	.980 $\pm$.020	.919 $\pm$.008	.113 $\pm$.100	.913 $\pm$.009
VAEAC	.846 $\pm$.003	.623 $\pm$.032	.900 $\pm$.000	.847 $\pm$.004	.655 $\pm$.027	.850 $\pm$.004
MIWAE	.871 $\pm$.008	.745 $\pm$.060	.980 $\pm$.020	.888 $\pm$.008	-.036 $\pm$.089	.901 $\pm$.008
not-MIWAE	.867 $\pm$.007	.596 $\pm$.073	1.000 $\pm$.000	.883 $\pm$.009	.017 $\pm$.139	.901 $\pm$.010
MIRACLE	.902 $\pm$.007	.525 $\pm$.069	1.000 $\pm$.000	.783 $\pm$.022	-.096 $\pm$.093	.832 $\pm$.020
ReMasker	.911 $\pm$.006	.497 $\pm$.044	1.000 $\pm$.000	.905 $\pm$.008	.405 $\pm$.125	.849 $\pm$.019
DrIM _{BASE}	.879 $\pm$.004	.723 $\pm$.083	1.000 $\pm$.000	.887 $\pm$.008	.171 $\pm$.166	.824 $\pm$.012
DrIM _{FINE}	.867 $\pm$.007	.674 $\pm$.075	.980 $\pm$.020	.932 $\pm$.004	.556 $\pm$.138	.891 $\pm$.011

model	concrete			kings		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.326 $\pm$.009	.175 $\pm$.126	.945 $\pm$.007	.440 $\pm$.005	.507 $\pm$.053	.978 $\pm$.002
kNNI	.349 $\pm$.007	.346 $\pm$.156	.940 $\pm$.013	.460 $\pm$.004	.410 $\pm$.064	.972 $\pm$.001
EM	.372 $\pm$.006	.419 $\pm$.210	.952 $\pm$.014	.459 $\pm$.003	.270 $\pm$.047	.945 $\pm$.002
SoftImpute	.328 $\pm$.008	.209 $\pm$.187	.945 $\pm$.010	.445 $\pm$.004	.310 $\pm$.010	.941 $\pm$.004
MICE	.345 $\pm$.008	.341 $\pm$.147	.960 $\pm$.009	.455 $\pm$.004	.290 $\pm$.043	.962 $\pm$.001
missForest	.340 $\pm$.006	.304 $\pm$.130	.952 $\pm$.012	.449 $\pm$.005	.470 $\pm$.058	.946 $\pm$.002
Sinkhorn	.341 $\pm$.006	.513 $\pm$.115	.933 $\pm$.021	.450 $\pm$.004	.490 $\pm$.059	.976 $\pm$.002
GAIN	.430 $\pm$.011	.491 $\pm$.074	.929 $\pm$.011	.491 $\pm$.004	.810 $\pm$.035	.952 $\pm$.004
VAEAC	.363 $\pm$.007	.553 $\pm$.062	.771 $\pm$.015	.513 $\pm$.002	.750 $\pm$.017	.862 $\pm$.002
MIWAE	.294 $\pm$.010	-.095 $\pm$.210	.926 $\pm$.018	.437 $\pm$.005	.450 $\pm$.056	.960 $\pm$.004
not-MIWAE	.313 $\pm$.009	.137 $\pm$.186	.955 $\pm$.014	.441 $\pm$.003	.470 $\pm$.047	.970 $\pm$.003
MIRACLE	.283 $\pm$.011	-.214 $\pm$.177	.924 $\pm$.014	.463 $\pm$.004	.280 $\pm$.055	.946 $\pm$.001
ReMasker	.269 $\pm$.014	.111 $\pm$.111	.921 $\pm$.015	.437 $\pm$.012	.390 $\pm$.108	.929 $\pm$.006
DrIM _{BASE}	.407 $\pm$.006	.426 $\pm$.108	.931 $\pm$.011	.498 $\pm$.008	.910 $\pm$.010	.964 $\pm$.004
DrIM _{FINE}	.397 $\pm$.010	.504 $\pm$.098	.936 $\pm$.022	.540 $\pm$.004	.950 $\pm$.017	.962 $\pm$.003
model	letter			loan		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.584 $\pm$.002	.700 $\pm$.000	.500 $\pm$.015	.916 $\pm$.001	.707 $\pm$.033	.895 $\pm$.007
kNNI	.674 $\pm$.002	.800 $\pm$.000	.948 $\pm$.002	.918 $\pm$.001	.783 $\pm$.047	.902 $\pm$.008
EM	.656 $\pm$.002	.800 $\pm$.000	.955 $\pm$.003	.919 $\pm$.001	.697 $\pm$.037	.878 $\pm$.017
SoftImpute	.594 $\pm$.002	.690 $\pm$.038	.668 $\pm$.015	.915 $\pm$.001	.724 $\pm$.046	.874 $\pm$.008
MICE	.621 $\pm$.002	.860 $\pm$.027	.920 $\pm$.005	.918 $\pm$.001	.733 $\pm$.047	.865 $\pm$.014
missForest	.614 $\pm$.001	.660 $\pm$.031	.742 $\pm$.016	.917 $\pm$.001	.787 $\pm$.054	.858 $\pm$.016
Sinkhorn	.625 $\pm$.001	.620 $\pm$.033	.789 $\pm$.011	.917 $\pm$.001	.646 $\pm$.043	.880 $\pm$.007
GAIN	.599 $\pm$.007	.640 $\pm$.033	.575 $\pm$.022	.915 $\pm$.002	.597 $\pm$.039	.846 $\pm$.010
VAEAC	.712 $\pm$.002	.780 $\pm$.013	.884 $\pm$.002	.827 $\pm$.002	.652 $\pm$.040	.849 $\pm$.011
MIWAE	.573 $\pm$.003	.812 $\pm$.031	.705 $\pm$.017	.914 $\pm$.001	.720 $\pm$.033	.872 $\pm$.007
not-MIWAE	.599 $\pm$.002	.620 $\pm$.013	.689 $\pm$.013	.916 $\pm$.000	.692 $\pm$.032	.876 $\pm$.010
MIRACLE	.648 $\pm$.006	.790 $\pm$.023	.940 $\pm$.004	.919 $\pm$.001	.730 $\pm$.054	.882 $\pm$.014
ReMasker	.645 $\pm$.002	.740 $\pm$.031	.931 $\pm$.004	.918 $\pm$.001	.727 $\pm$.040	.887 $\pm$.016
DrIM _{BASE}	.661 $\pm$.003	.790 $\pm$.023	.952 $\pm$.003	.910 $\pm$.001	.647 $\pm$.030	.878 $\pm$.006
DrIM _{FINE}	.716 $\pm$.004	1.000 $\pm$.000	.967 $\pm$.002	.912 $\pm$.001	.690 $\pm$.038	.874 $\pm$.006
model	redwine			whitewine		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.426 $\pm$.007	.206 $\pm$.090	.905 $\pm$.014	.373 $\pm$.005	.086 $\pm$.122	.593 $\pm$.048
kNNI	.440 $\pm$.006	.171 $\pm$.137	.909 $\pm$.021	.386 $\pm$.004	.150 $\pm$.079	.917 $\pm$.008
EM	.435 $\pm$.006	.102 $\pm$.144	.892 $\pm$.017	.384 $\pm$.004	.260 $\pm$.056	.809 $\pm$.023
SoftImpute	.430 $\pm$.007	.047 $\pm$.124	.869 $\pm$.018	.377 $\pm$.005	.160 $\pm$.083	.906 $\pm$.026
MICE	.428 $\pm$.006	-.022 $\pm$.133	.895 $\pm$.020	.388 $\pm$.005	.190 $\pm$.066	.921 $\pm$.021
missForest	.434 $\pm$.005	.001 $\pm$.131	.859 $\pm$.023	.383 $\pm$.004	.070 $\pm$.088	.871 $\pm$.009
Sinkhorn	.435 $\pm$.005	.044 $\pm$.123	.884 $\pm$.018	.386 $\pm$.005	.120 $\pm$.059	.894 $\pm$.013
GAIN	.509 $\pm$.007	.404 $\pm$.110	.897 $\pm$.014	.425 $\pm$.007	.630 $\pm$.072	.846 $\pm$.028
VAEAC	.483 $\pm$.005	.263 $\pm$.046	.785 $\pm$.019	.423 $\pm$.004	.450 $\pm$.017	.722 $\pm$.025
MIWAE	.444 $\pm$.007	-.038 $\pm$.140	.883 $\pm$.017	.380 $\pm$.004	.240 $\pm$.090	.925 $\pm$.007
not-MIWAE	.432 $\pm$.007	.061 $\pm$.151	.909 $\pm$.011	.384 $\pm$.004	.150 $\pm$.092	.914 $\pm$.022
MIRACLE	.430 $\pm$.007	.001 $\pm$.085	.877 $\pm$.020	.379 $\pm$.006	.290 $\pm$.050	.855 $\pm$.018
ReMasker	.415 $\pm$.013	-.030 $\pm$.099	.785 $\pm$.028	.370 $\pm$.008	-.050 $\pm$.120	.563 $\pm$.048
DrIM _{BASE}	.532 $\pm$.006	.353 $\pm$.110	.911 $\pm$.014	.476 $\pm$.005	.321 $\pm$.057	.902 $\pm$.017
DrIM _{FINE}	.536 $\pm$.008	.448 $\pm$.062	.892 $\pm$.016	.487 $\pm$.006	.417 $\pm$.047	.840 $\pm$.020

Table 14: Machine learning utility for each dataset under **MAR** at 0.3 missingness. The means and the standard errors of the mean across 10 repeated experiments are reported. $\uparrow$ denotes higher is better.

model	abalone			anuran		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.163 $\pm$.014	.147 $\pm$.163	.836 $\pm$.021	.830 $\pm$.018	-.046 $\pm$.150	.552 $\pm$.059
kNNI	.167 $\pm$.013	.142 $\pm$.177	.938 $\pm$.016	.816 $\pm$.019	-.061 $\pm$.149	.683 $\pm$.040
EM	.164 $\pm$.013	.182 $\pm$.140	.902 $\pm$.037	.825 $\pm$.018	-.152 $\pm$.159	.658 $\pm$.044
SoftImpute	.161 $\pm$.012	.211 $\pm$.133	.926 $\pm$.010	.827 $\pm$.018	.005 $\pm$.134	.671 $\pm$.050
MICE	.161 $\pm$.013	.283 $\pm$.149	.962 $\pm$.012	.831 $\pm$.017	-.051 $\pm$.145	.699 $\pm$.046
missForest	.161 $\pm$.013	.313 $\pm$.159	.869 $\pm$.041	.831 $\pm$.017	-.015 $\pm$.157	.650 $\pm$.043
Sinkhorn	.164 $\pm$.013	.044 $\pm$.202	.960 $\pm$.010	.825 $\pm$.018	.055 $\pm$.131	.668 $\pm$.045
GAIN	.206 $\pm$.006	.415 $\pm$.171	.905 $\pm$.015	.863 $\pm$.023	.451 $\pm$.066	.808 $\pm$.051
VAEAC	.131 $\pm$.005	.497 $\pm$.043	.824 $\pm$.015	.891 $\pm$.001	.565 $\pm$.076	.871 $\pm$.003
MIWAE	.160 $\pm$.013	-.013 $\pm$.194	.931 $\pm$.010	.825 $\pm$.018	-.096 $\pm$.162	.631 $\pm$.045
not-MIWAE	.164 $\pm$.013	.055 $\pm$.157	.914 $\pm$.019	.825 $\pm$.018	-.065 $\pm$.147	.681 $\pm$.043
MIRACLE	.163 $\pm$.013	.146 $\pm$.146	.900 $\pm$.035	.822 $\pm$.018	-.113 $\pm$.144	.701 $\pm$.046
ReMasker	.169 $\pm$.012	.283 $\pm$.173	.888 $\pm$.043	.815 $\pm$.019	-.183 $\pm$.170	.668 $\pm$.041
DrIM _{BASE}	.202 $\pm$.007	.688 $\pm$.058	.917 $\pm$.017	.897 $\pm$.012	-.188 $\pm$.197	.904 $\pm$.012
DrIM _{FINE}	.197 $\pm$.009	.606 $\pm$.118	.879 $\pm$.025	.971 $\pm$.005	.210 $\pm$.125	.929 $\pm$.012

model	banknote			breast		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.861 $\pm$.020	.786 $\pm$.065	.800 $\pm$.073	.921 $\pm$.013	.586 $\pm$.098	.810 $\pm$.038
kNNI	.876 $\pm$.015	.633 $\pm$.074	.940 $\pm$.031	.911 $\pm$.019	.462 $\pm$.124	.898 $\pm$.021
EM	.884 $\pm$.015	.726 $\pm$.058	.940 $\pm$.031	.909 $\pm$.020	.432 $\pm$.141	.861 $\pm$.024
SoftImpute	.868 $\pm$.020	.673 $\pm$.053	.960 $\pm$.027	.912 $\pm$.018	.536 $\pm$.122	.903 $\pm$.021
MICE	.881 $\pm$.020	.673 $\pm$.078	.940 $\pm$.031	.910 $\pm$.021	.475 $\pm$.182	.912 $\pm$.018
missForest	.889 $\pm$.017	.623 $\pm$.093	.960 $\pm$.027	.910 $\pm$.019	.375 $\pm$.181	.910 $\pm$.020
Sinkhorn	.873 $\pm$.020	.658 $\pm$.066	.960 $\pm$.027	.917 $\pm$.015	.497 $\pm$.125	.905 $\pm$.020
GAIN	.934 $\pm$.005	.729 $\pm$.044	1.000 $\pm$.000	.923 $\pm$.013	.674 $\pm$.097	.914 $\pm$.018
VAEAC	.849 $\pm$.004	.678 $\pm$.036	.860 $\pm$.027	.848 $\pm$.003	.622 $\pm$.042	.853 $\pm$.005
MIWAE	.865 $\pm$.020	.840 $\pm$.069	.900 $\pm$.033	.908 $\pm$.022	.436 $\pm$.163	.904 $\pm$.023
not-MIWAE	.863 $\pm$.020	.570 $\pm$.097	.880 $\pm$.033	.909 $\pm$.019	.550 $\pm$.150	.895 $\pm$.024
MIRACLE	.889 $\pm$.015	.709 $\pm$.055	.960 $\pm$.027	.883 $\pm$.021	.250 $\pm$.169	.770 $\pm$.027
ReMasker	.887 $\pm$.018	.740 $\pm$.045	.940 $\pm$.031	.905 $\pm$.017	.316 $\pm$.183	.816 $\pm$.022
DrIM _{BASE}	.906 $\pm$.008	.776 $\pm$.072	.960 $\pm$.027	.877 $\pm$.023	.485 $\pm$.157	.777 $\pm$.057
DrIM _{FINE}	.879 $\pm$.014	.711 $\pm$.081	.940 $\pm$.031	.931 $\pm$.008	.672 $\pm$.085	.893 $\pm$.018

model	concrete			kings		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.350 $\pm$.025	.511 $\pm$.100	.881 $\pm$.028	.421 $\pm$.009	.678 $\pm$.083	.900 $\pm$.017
kNNI	.372 $\pm$.024	.769 $\pm$.098	.929 $\pm$.013	.432 $\pm$.009	.656 $\pm$.084	.957 $\pm$.005
EM	.391 $\pm$.024	.733 $\pm$.043	.845 $\pm$.037	.423 $\pm$.010	.689 $\pm$.065	.913 $\pm$.012
SoftImpute	.349 $\pm$.023	.620 $\pm$.101	.879 $\pm$.026	.419 $\pm$.010	.600 $\pm$.091	.938 $\pm$.006
MICE	.370 $\pm$.027	.549 $\pm$.132	.886 $\pm$.029	.423 $\pm$.010	.562 $\pm$.089	.942 $\pm$.007
missForest	.353 $\pm$.022	.601 $\pm$.149	.826 $\pm$.033	.428 $\pm$.010	.700 $\pm$.087	.948 $\pm$.005
Sinkhorn	.370 $\pm$.023	.595 $\pm$.164	.860 $\pm$.020	.423 $\pm$.009	.725 $\pm$.070	.966 $\pm$.004
GAIN	.421 $\pm$.013	.624 $\pm$.080	.869 $\pm$.022	.500 $\pm$.011	.760 $\pm$.048	.959 $\pm$.008
VAEAC	.357 $\pm$.007	.649 $\pm$.051	.721 $\pm$.027	.514 $\pm$.001	.770 $\pm$.015	.860 $\pm$.006
MIWAE	.335 $\pm$.023	.389 $\pm$.111	.845 $\pm$.038	.416 $\pm$.010	.644 $\pm$.078	.949 $\pm$.006
not-MIWAE	.338 $\pm$.023	.602 $\pm$.089	.879 $\pm$.026	.423 $\pm$.008	.690 $\pm$.078	.949 $\pm$.006
MIRACLE	.339 $\pm$.025	.247 $\pm$.147	.893 $\pm$.019	.436 $\pm$.009	.578 $\pm$.074	.950 $\pm$.010
ReMasker	.332 $\pm$.025	.263 $\pm$.084	.893 $\pm$.019	.423 $\pm$.014	.533 $\pm$.078	.935 $\pm$.006
DrIM _{BASE}	.415 $\pm$.011	.666 $\pm$.080	.900 $\pm$.019	.512 $\pm$.008	.890 $\pm$.035	.948 $\pm$.009
DrIM _{FINE}	.409 $\pm$.009	.690 $\pm$.069	.879 $\pm$.025	.548 $\pm$.009	.960 $\pm$.016	.966 $\pm$.003

model	letter			loan		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.621 $\pm$.028	.840 $\pm$.050	.715 $\pm$.032	.915 $\pm$.003	.670 $\pm$.063	.898 $\pm$.020
kNNI	.664 $\pm$.028	.832 $\pm$.029	.936 $\pm$.012	.919 $\pm$.002	.804 $\pm$.047	.935 $\pm$.019
EM	.666 $\pm$.025	.840 $\pm$.027	.952 $\pm$.006	.893 $\pm$.022	.732 $\pm$.048	.910 $\pm$.021
SoftImpute	.627 $\pm$.029	.880 $\pm$.025	.781 $\pm$.023	.916 $\pm$.003	.673 $\pm$.042	.910 $\pm$.021
MICE	.646 $\pm$.028	.890 $\pm$.023	.894 $\pm$.014	.917 $\pm$.002	.759 $\pm$.052	.908 $\pm$.027
missForest	.637 $\pm$.027	.857 $\pm$.040	.803 $\pm$.027	.918 $\pm$.002	.761 $\pm$.045	.902 $\pm$.026
Sinkhorn	.646 $\pm$.027	.840 $\pm$.045	.843 $\pm$.020	.915 $\pm$.003	.731 $\pm$.039	.912 $\pm$.021
GAIN	.645 $\pm$.027	.810 $\pm$.038	.794 $\pm$.017	.915 $\pm$.003	.572 $\pm$.059	.888 $\pm$.020
VAEAC	.724 $\pm$.002	.800 $\pm$.000	.890 $\pm$.002	.827 $\pm$.001	.640 $\pm$.054	.821 $\pm$.020
MIWAE	.616 $\pm$.029	.880 $\pm$.025	.809 $\pm$.010	.916 $\pm$.003	.802 $\pm$.052	.891 $\pm$.023
not-MIWAE	.629 $\pm$.028	.897 $\pm$.021	.805 $\pm$.020	.916 $\pm$.003	.744 $\pm$.048	.908 $\pm$.021
MIRACLE	.671 $\pm$.027	.840 $\pm$.027	.951 $\pm$.009	.920 $\pm$.002	.834 $\pm$.040	.931 $\pm$.025
ReMasker	.659 $\pm$.027	.840 $\pm$.027	.947 $\pm$.008	.919 $\pm$.002	.803 $\pm$.048	.910 $\pm$.026
DrIM _{BASE}	.659 $\pm$.025	.780 $\pm$.047	.947 $\pm$.006	.915 $\pm$.003	.730 $\pm$.040	.902 $\pm$.021
DrIM _{FINE}	.691 $\pm$.020	.940 $\pm$.022	.954 $\pm$.007	.915 $\pm$.003	.757 $\pm$.041	.880 $\pm$.019

model	redwine			whitewine		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.428 $\pm$.027	.673 $\pm$.092	.543 $\pm$.085	.386 $\pm$.021	.647 $\pm$.096	.525 $\pm$.071
kNNI	.437 $\pm$.027	.697 $\pm$.115	.716 $\pm$.064	.395 $\pm$.021	.675 $\pm$.044	.757 $\pm$.056
EM	.436 $\pm$.024	.502 $\pm$.122	.699 $\pm$.053	.398 $\pm$.022	.660 $\pm$.096	.667 $\pm$.078
SoftImpute	.421 $\pm$.026	.571 $\pm$.116	.689 $\pm$.067	.387 $\pm$.023	.777 $\pm$.078	.806 $\pm$.038
MICE	.430 $\pm$.026	.633 $\pm$.103	.686 $\pm$.058	.390 $\pm$.024	.660 $\pm$.081	.769 $\pm$.054
missForest	.426 $\pm$.026	.597 $\pm$.108	.717 $\pm$.061	.387 $\pm$.023	.687 $\pm$.093	.753 $\pm$.049
Sinkhorn	.431 $\pm$.025	.679 $\pm$.109	.686 $\pm$.056	.389 $\pm$.022	.660 $\pm$.086	.714 $\pm$.043
GAIN	.526 $\pm$.009	.382 $\pm$.078	.812 $\pm$.028	.467 $\pm$.013	.600 $\pm$.074	.822 $\pm$.062
VAEAC	.485 $\pm$.004	.403 $\pm$.053	.808 $\pm$.018	.431 $\pm$.002	.500 $\pm$.037	.757 $\pm$.025
MIWAE	.431 $\pm$.026	.602 $\pm$.094	.676 $\pm$.057	.389 $\pm$.023	.720 $\pm$.074	.814 $\pm$.035
not-MIWAE	.420 $\pm$.027	.622 $\pm$.118	.704 $\pm$.057	.387 $\pm$.022	.740 $\pm$.048	.773 $\pm$.044
MIRACLE	.435 $\pm$.024	.487 $\pm$.125	.731 $\pm$.057	.393 $\pm$.022	.760 $\pm$.070	.731 $\pm$.045
ReMasker	.415 $\pm$.029	.603 $\pm$.092	.567 $\pm$.074	.381 $\pm$.022	.605 $\pm$.104	.615 $\pm$.086
DrIM _{BASE}	.539 $\pm$.010	.634 $\pm$.078	.853 $\pm$.036	.471 $\pm$.009	.707 $\pm$.062	.835 $\pm$.030
DrIM _{FINE}	.555 $\pm$.009	.451 $\pm$.095	.878 $\pm$.015	.498 $\pm$.004	.580 $\pm$.066	.858 $\pm$.023

Table 15: Machine learning utility for each dataset under **MNARL** at 0.3 missingness. The means and the standard errors of the mean across 10 repeated experiments are reported. $\uparrow$ denotes higher is better.

model	abalone			anuran		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.131 $\pm$.005	-.206 $\pm$.080	.895 $\pm$.018	.817 $\pm$.010	.035 $\pm$.100	.683 $\pm$.039
kNNI	.140 $\pm$.004	-.220 $\pm$.133	.924 $\pm$.017	.788 $\pm$.007	-.233 $\pm$.126	.513 $\pm$.029
EM	.133 $\pm$.005	-.238 $\pm$.125	.926 $\pm$.013	.801 $\pm$.006	-.313 $\pm$.113	.519 $\pm$.038
SoftImpute	.125 $\pm$.005	-.297 $\pm$.137	.914 $\pm$.016	.811 $\pm$.007	-.090 $\pm$.123	.576 $\pm$.047
MICE	.126 $\pm$.005	-.457 $\pm$.070	.955 $\pm$.013	.813 $\pm$.007	-.237 $\pm$.131	.577 $\pm$.044
missForest	.129 $\pm$.005	-.282 $\pm$.104	.893 $\pm$.023	.810 $\pm$.007	-.166 $\pm$.139	.518 $\pm$.032
Sinkhorn	.133 $\pm$.004	-.374 $\pm$.093	.940 $\pm$.010	.803 $\pm$.008	-.172 $\pm$.125	.505 $\pm$.032
GAIN	.187 $\pm$.005	-.008 $\pm$.180	.898 $\pm$.027	.875 $\pm$.020	.088 $\pm$.117	.751 $\pm$.043
VAEAC	.109 $\pm$.023	.369 $\pm$.081	.740 $\pm$.095	.890 $\pm$.001	.538 $\pm$.081	.870 $\pm$.004
MIWAE	.121 $\pm$.005	-.417 $\pm$.074	.921 $\pm$.017	.811 $\pm$.007	-.133 $\pm$.142	.522 $\pm$.032
not-MIWAE	.126 $\pm$.005	-.387 $\pm$.100	.933 $\pm$.019	.809 $\pm$.008	-.141 $\pm$.156	.547 $\pm$.038
MIRACLE	.129 $\pm$.005	-.276 $\pm$.076	.914 $\pm$.012	.787 $\pm$.009	-.304 $\pm$.104	.589 $\pm$.035
ReMasker	.122 $\pm$.006	-.417 $\pm$.090	.910 $\pm$.025	.789 $\pm$.007	-.283 $\pm$.135	.494 $\pm$.036
DrIM _{BASE}	.196 $\pm$.004	.639 $\pm$.074	.940 $\pm$.011	.891 $\pm$.007	-.272 $\pm$.168	.907 $\pm$.007
DrIM _{FINE}	.191 $\pm$.007	.394 $\pm$.138	.890 $\pm$.029	.960 $\pm$.005	-.044 $\pm$.117	.933 $\pm$.005

model	banknote			breast		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.839 $\pm$.015	.677 $\pm$.084	.860 $\pm$.060	.903 $\pm$.009	.529 $\pm$.116	.783 $\pm$.027
kNNI	.846 $\pm$.014	.701 $\pm$.089	.940 $\pm$.031	.890 $\pm$.014	.175 $\pm$.115	.866 $\pm$.021
EM	.886 $\pm$.013	.690 $\pm$.040	.940 $\pm$.031	.889 $\pm$.017	.047 $\pm$.127	.861 $\pm$.023
SoftImpute	.859 $\pm$.014	.503 $\pm$.060	.920 $\pm$.033	.890 $\pm$.015	.185 $\pm$.123	.849 $\pm$.022
MICE	.879 $\pm$.016	.627 $\pm$.060	.940 $\pm$.031	.889 $\pm$.018	.161 $\pm$.139	.865 $\pm$.024
missForest	.879 $\pm$.014	.583 $\pm$.093	.940 $\pm$.031	.889 $\pm$.016	.119 $\pm$.156	.861 $\pm$.026
Sinkhorn	.857 $\pm$.014	.658 $\pm$.073	.940 $\pm$.031	.898 $\pm$.012	.201 $\pm$.138	.871 $\pm$.022
GAIN	.901 $\pm$.011	.587 $\pm$.059	.980 $\pm$.020	.922 $\pm$.006	.525 $\pm$.158	.865 $\pm$.031
VAEAC	.845 $\pm$.003	.612 $\pm$.047	.900 $\pm$.000	.843 $\pm$.003	.644 $\pm$.047	.841 $\pm$.004
MIWAE	.848 $\pm$.016	.713 $\pm$.077	.940 $\pm$.031	.889 $\pm$.018	.088 $\pm$.165	.872 $\pm$.020
not-MIWAE	.840 $\pm$.014	.685 $\pm$.046	1.000 $\pm$.000	.891 $\pm$.017	.162 $\pm$.149	.879 $\pm$.020
MIRACLE	.878 $\pm$.014	.577 $\pm$.050	.920 $\pm$.061	.818 $\pm$.017	-.140 $\pm$.162	.755 $\pm$.033
ReMasker	.874 $\pm$.014	.580 $\pm$.094	.960 $\pm$.027	.901 $\pm$.007	.274 $\pm$.113	.759 $\pm$.035
DrIM _{BASE}	.877 $\pm$.010	.767 $\pm$.060	.920 $\pm$.033	.858 $\pm$.018	.293 $\pm$.131	.737 $\pm$.034
DrIM _{FINE}	.877 $\pm$.008	.704 $\pm$.091	.960 $\pm$.027	.928 $\pm$.007	.727 $\pm$.082	.884 $\pm$.013

model	concrete			kings		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.307 $\pm$.013	.184 $\pm$.110	.917 $\pm$.020	.413 $\pm$.010	.730 $\pm$.065	.959 $\pm$.005
kNNI	.328 $\pm$.009	.534 $\pm$.093	.869 $\pm$.028	.432 $\pm$.010	.640 $\pm$.095	.956 $\pm$.003
EM	.359 $\pm$.013	.654 $\pm$.080	.707 $\pm$.070	.432 $\pm$.009	.540 $\pm$.083	.926 $\pm$.010
SoftImpute	.307 $\pm$.013	.184 $\pm$.110	.917 $\pm$.020	.416 $\pm$.009	.500 $\pm$.099	.926 $\pm$.004
MICE	.322 $\pm$.012	.433 $\pm$.104	.740 $\pm$.075	.430 $\pm$.009	.390 $\pm$.075	.942 $\pm$.004
missForest	.317 $\pm$.010	.227 $\pm$.090	.788 $\pm$.055	.423 $\pm$.009	.680 $\pm$.074	.936 $\pm$.004
Sinkhorn	.324 $\pm$.010	.393 $\pm$.139	.762 $\pm$.061	.422 $\pm$.008	.730 $\pm$.065	.962 $\pm$.003
GAIN	.414 $\pm$.009	.371 $\pm$.153	.900 $\pm$.017	.488 $\pm$.008	.770 $\pm$.037	.963 $\pm$.007
VAEAC	.347 $\pm$.010	.472 $\pm$.056	.743 $\pm$.036	.510 $\pm$.002	.760 $\pm$.016	.854 $\pm$.002
MIWAE	.276 $\pm$.012	.095 $\pm$.167	.843 $\pm$.041	.415 $\pm$.008	.730 $\pm$.079	.949 $\pm$.002
not-MIWAE	.290 $\pm$.010	.441 $\pm$.145	.831 $\pm$.041	.410 $\pm$.010	.682 $\pm$.072	.947 $\pm$.003
MIRACLE	.289 $\pm$.019	-.064 $\pm$.121	.860 $\pm$.032	.437 $\pm$.007	.460 $\pm$.095	.941 $\pm$.005
ReMasker	.273 $\pm$.018	.141 $\pm$.141	.905 $\pm$.019	.428 $\pm$.012	.430 $\pm$.092	.905 $\pm$.011
DrIM _{BASE}	.401 $\pm$.007	.386 $\pm$.144	.924 $\pm$.016	.490 $\pm$.010	.920 $\pm$.029	.955 $\pm$.004
DrIM _{FINE}	.408 $\pm$.008	.489 $\pm$.091	.929 $\pm$.014	.530 $\pm$.007	.940 $\pm$.016	.960 $\pm$.002
model	letter			loan		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.561 $\pm$.005	.740 $\pm$.048	.665 $\pm$.030	.912 $\pm$.002	.630 $\pm$.065	.881 $\pm$.023
kNNI	.651 $\pm$.007	.800 $\pm$.000	.950 $\pm$.005	.916 $\pm$.002	.740 $\pm$.027	.902 $\pm$.020
EM	.632 $\pm$.004	.810 $\pm$.028	.953 $\pm$.005	.916 $\pm$.001	.697 $\pm$.038	.855 $\pm$.020
SoftImpute	.569 $\pm$.005	.777 $\pm$.051	.749 $\pm$.026	.913 $\pm$.002	.702 $\pm$.040	.868 $\pm$.014
MICE	.595 $\pm$.006	.840 $\pm$.027	.906 $\pm$.012	.917 $\pm$.002	.714 $\pm$.040	.859 $\pm$.021
missForest	.586 $\pm$.005	.740 $\pm$.052	.796 $\pm$.022	.914 $\pm$.002	.750 $\pm$.052	.876 $\pm$.018
Sinkhorn	.599 $\pm$.006	.730 $\pm$.050	.856 $\pm$.016	.912 $\pm$.003	.688 $\pm$.052	.884 $\pm$.019
GAIN	.580 $\pm$.010	.620 $\pm$.042	.690 $\pm$.037	.911 $\pm$.002	.587 $\pm$.015	.890 $\pm$.021
VAEAC	.712 $\pm$.002	.800 $\pm$.000	.887 $\pm$.002	.825 $\pm$.001	.620 $\pm$.042	.838 $\pm$.014
MIWAE	.549 $\pm$.005	.860 $\pm$.022	.772 $\pm$.025	.914 $\pm$.002	.713 $\pm$.052	.858 $\pm$.022
not-MIWAE	.573 $\pm$.005	.740 $\pm$.050	.760 $\pm$.028	.912 $\pm$.002	.710 $\pm$.035	.869 $\pm$.019
MIRACLE	.620 $\pm$.009	.820 $\pm$.013	.932 $\pm$.007	.918 $\pm$.002	.783 $\pm$.040	.878 $\pm$.022
ReMasker	.624 $\pm$.005	.780 $\pm$.020	.942 $\pm$.007	.918 $\pm$.002	.760 $\pm$.041	.882 $\pm$.024
DrIM _{BASE}	.631 $\pm$.004	.760 $\pm$.027	.940 $\pm$.005	.911 $\pm$.003	.626 $\pm$.045	.862 $\pm$.015
DrIM _{FINE}	.691 $\pm$.008	.950 $\pm$.027	.957 $\pm$.006	.910 $\pm$.002	.720 $\pm$.042	.863 $\pm$.014
model	redwine			whitewine		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.375 $\pm$.009	.422 $\pm$.114	.682 $\pm$.072	.344 $\pm$.007	.540 $\pm$.117	.648 $\pm$.069
kNNI	.400 $\pm$.010	.477 $\pm$.161	.705 $\pm$.049	.357 $\pm$.008	.577 $\pm$.119	.804 $\pm$.047
EM	.395 $\pm$.009	.529 $\pm$.121	.666 $\pm$.052	.360 $\pm$.008	.656 $\pm$.097	.688 $\pm$.070
SoftImpute	.386 $\pm$.010	.426 $\pm$.113	.671 $\pm$.050	.342 $\pm$.007	.530 $\pm$.125	.792 $\pm$.044
MICE	.390 $\pm$.012	.467 $\pm$.148	.677 $\pm$.052	.349 $\pm$.009	.450 $\pm$.138	.788 $\pm$.062
missForest	.390 $\pm$.011	.402 $\pm$.142	.712 $\pm$.050	.350 $\pm$.007	.606 $\pm$.109	.773 $\pm$.055
Sinkhorn	.389 $\pm$.011	.454 $\pm$.135	.720 $\pm$.052	.352 $\pm$.006	.465 $\pm$.143	.774 $\pm$.059
GAIN	.506 $\pm$.007	.382 $\pm$.110	.796 $\pm$.031	.427 $\pm$.010	.537 $\pm$.075	.695 $\pm$.095
VAEAC	.472 $\pm$.005	.266 $\pm$.115	.785 $\pm$.016	.423 $\pm$.004	.450 $\pm$.031	.700 $\pm$.021
MIWAE	.392 $\pm$.010	.506 $\pm$.121	.715 $\pm$.060	.347 $\pm$.008	.537 $\pm$.130	.817 $\pm$.046
not-MIWAE	.383 $\pm$.011	.479 $\pm$.116	.662 $\pm$.055	.346 $\pm$.007	.575 $\pm$.115	.755 $\pm$.050
MIRACLE	.399 $\pm$.009	.309 $\pm$.145	.699 $\pm$.056	.360 $\pm$.005	.470 $\pm$.114	.765 $\pm$.058
ReMasker	.388 $\pm$.020	.185 $\pm$.098	.534 $\pm$.056	.339 $\pm$.010	.342 $\pm$.169	.532 $\pm$.076
DrIM _{BASE}	.525 $\pm$.008	.334 $\pm$.098	.839 $\pm$.021	.470 $\pm$.006	.510 $\pm$.072	.837 $\pm$.032
DrIM _{FINE}	.544 $\pm$.005	.105 $\pm$.072	.902 $\pm$.013	.490 $\pm$.005	.541 $\pm$.080	.872 $\pm$.027

Table 16: Machine learning utility for each dataset under **MNARQ** at 0.3 missingness.. The means and the standard errors of the mean across 10 repeated experiments are reported. $\uparrow$ denotes higher is better.

model	abalone			anuran		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.143 $\pm$.015	-.292 $\pm$.211	.774 $\pm$.035	.921 $\pm$.021	.278 $\pm$.118	.549 $\pm$.075
kNNI	.146 $\pm$.013	-.284 $\pm$.198	.926 $\pm$.021	.914 $\pm$.023	.333 $\pm$.128	.700 $\pm$.052
EM	.142 $\pm$.013	-.317 $\pm$.173	.855 $\pm$.033	.911 $\pm$.021	.223 $\pm$.120	.648 $\pm$.052
SoftImpute	.141 $\pm$.014	-.262 $\pm$.199	.924 $\pm$.023	.918 $\pm$.022	.300 $\pm$.118	.717 $\pm$.063
MICE	.143 $\pm$.015	-.213 $\pm$.194	.950 $\pm$.017	.920 $\pm$.021	.344 $\pm$.120	.738 $\pm$.052
missForest	.145 $\pm$.014	-.262 $\pm$.205	.919 $\pm$.024	.919 $\pm$.022	.325 $\pm$.125	.687 $\pm$.050
Sinkhorn	.143 $\pm$.013	-.154 $\pm$.172	.929 $\pm$.020	.917 $\pm$.022	.316 $\pm$.121	.702 $\pm$.053
GAIN	.182 $\pm$.012	.021 $\pm$.198	.907 $\pm$.019	.951 $\pm$.014	.490 $\pm$.100	.831 $\pm$.028
VAEAC	.136 $\pm$.004	.538 $\pm$.049	.829 $\pm$.011	.893 $\pm$.001	.613 $\pm$.082	.871 $\pm$.003
MIWAE	.138 $\pm$.016	-.307 $\pm$.202	.948 $\pm$.010	.921 $\pm$.021	.364 $\pm$.128	.700 $\pm$.058
not-MIWAE	.141 $\pm$.014	-.217 $\pm$.197	.950 $\pm$.015	.919 $\pm$.022	.356 $\pm$.122	.726 $\pm$.053
MIRACLE	.146 $\pm$.014	-.224 $\pm$.185	.924 $\pm$.025	.917 $\pm$.022	.288 $\pm$.126	.726 $\pm$.052
ReMasker	.148 $\pm$.013	-.197 $\pm$.154	.798 $\pm$.037	.913 $\pm$.024	.267 $\pm$.112	.697 $\pm$.050
DrIM _{BASE}	.209 $\pm$.005	.827 $\pm$.051	.876 $\pm$.022	.953 $\pm$.012	.631 $\pm$.069	.936 $\pm$.013
DrIM _{FINE}	.202 $\pm$.006	.524 $\pm$.126	.817 $\pm$.028	.975 $\pm$.007	.571 $\pm$.162	.928 $\pm$.010

model	banknote			breast		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.904 $\pm$.012	.832 $\pm$.053	.800 $\pm$.052	.915 $\pm$.011	.469 $\pm$.164	.841 $\pm$.032
kNNI	.916 $\pm$.011	.680 $\pm$.075	1.000 $\pm$.000	.915 $\pm$.011	.513 $\pm$.176	.923 $\pm$.016
EM	.922 $\pm$.007	.753 $\pm$.071	.960 $\pm$.027	.915 $\pm$.012	.333 $\pm$.205	.884 $\pm$.019
SoftImpute	.911 $\pm$.011	.604 $\pm$.087	1.000 $\pm$.000	.915 $\pm$.010	.406 $\pm$.202	.915 $\pm$.018
MICE	.922 $\pm$.010	.711 $\pm$.070	1.000 $\pm$.000	.918 $\pm$.013	.231 $\pm$.216	.906 $\pm$.025
missForest	.920 $\pm$.010	.600 $\pm$.066	1.000 $\pm$.000	.915 $\pm$.012	.380 $\pm$.175	.927 $\pm$.013
Sinkhorn	.912 $\pm$.011	.670 $\pm$.073	1.000 $\pm$.000	.920 $\pm$.011	.409 $\pm$.195	.919 $\pm$.022
GAIN	.943 $\pm$.004	.720 $\pm$.043	1.000 $\pm$.000	.929 $\pm$.009	.424 $\pm$.161	.925 $\pm$.020
VAEAC	.854 $\pm$.003	.625 $\pm$.035	.900 $\pm$.000	.850 $\pm$.003	.624 $\pm$.044	.850 $\pm$.005
MIWAE	.904 $\pm$.013	.774 $\pm$.050	1.000 $\pm$.000	.919 $\pm$.011	.412 $\pm$.195	.900 $\pm$.027
not-MIWAE	.902 $\pm$.012	.819 $\pm$.045	.980 $\pm$.020	.915 $\pm$.012	.474 $\pm$.171	.912 $\pm$.021
MIRACLE	.919 $\pm$.013	.739 $\pm$.056	1.000 $\pm$.000	.894 $\pm$.017	.088 $\pm$.174	.788 $\pm$.038
ReMasker	.918 $\pm$.010	.663 $\pm$.076	1.000 $\pm$.000	.911 $\pm$.009	.194 $\pm$.181	.846 $\pm$.028
DrIM _{BASE}	.919 $\pm$.009	.770 $\pm$.048	1.000 $\pm$.000	.909 $\pm$.013	.366 $\pm$.155	.878 $\pm$.023
DrIM _{FINE}	.917 $\pm$.009	.753 $\pm$.056	1.000 $\pm$.000	.936 $\pm$.004	.764 $\pm$.062	.917 $\pm$.012

model	concrete			kings		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.351 $\pm$.021	.223 $\pm$.175	.945 $\pm$.009	.521 $\pm$.024	.610 $\pm$.119	.952 $\pm$.006
kNNI	.370 $\pm$.020	.634 $\pm$.091	.819 $\pm$.032	.535 $\pm$.025	.580 $\pm$.125	.976 $\pm$.001
EM	.381 $\pm$.019	.562 $\pm$.112	.660 $\pm$.067	.506 $\pm$.027	.529 $\pm$.122	.965 $\pm$.006
SoftImpute	.351 $\pm$.020	.386 $\pm$.123	.807 $\pm$.057	.525 $\pm$.025	.580 $\pm$.131	.962 $\pm$.003
MICE	.367 $\pm$.022	.418 $\pm$.152	.717 $\pm$.067	.532 $\pm$.024	.535 $\pm$.124	.961 $\pm$.006
missForest	.360 $\pm$.021	.344 $\pm$.134	.705 $\pm$.069	.526 $\pm$.024	.660 $\pm$.109	.963 $\pm$.005
Sinkhorn	.360 $\pm$.021	.639 $\pm$.094	.638 $\pm$.074	.529 $\pm$.024	.630 $\pm$.122	.982 $\pm$.001
GAIN	.435 $\pm$.012	.623 $\pm$.051	.731 $\pm$.072	.558 $\pm$.018	.830 $\pm$.038	.973 $\pm$.005
VAEAC	.370 $\pm$.011	.552 $\pm$.065	.557 $\pm$.069	.519 $\pm$.002	.760 $\pm$.016	.872 $\pm$.004
MIWAE	.330 $\pm$.021	.170 $\pm$.154	.871 $\pm$.029	.520 $\pm$.025	.600 $\pm$.134	.967 $\pm$.002
not-MIWAE	.350 $\pm$.020	.334 $\pm$.155	.838 $\pm$.040	.527 $\pm$.025	.630 $\pm$.122	.974 $\pm$.003
MIRACLE	.336 $\pm$.020	-.174 $\pm$.132	.829 $\pm$.044	.536 $\pm$.025	.550 $\pm$.123	.968 $\pm$.006
ReMasker	.309 $\pm$.021	.146 $\pm$.143	.886 $\pm$.020	.518 $\pm$.022	.620 $\pm$.119	.949 $\pm$.007
DrIM _{BASE}	.416 $\pm$.011	.550 $\pm$.118	.869 $\pm$.027	.563 $\pm$.012	.890 $\pm$.043	.976 $\pm$.003
DrIM _{FINE}	.416 $\pm$.012	.680 $\pm$.087	.907 $\pm$.015	.586 $\pm$.007	.947 $\pm$.016	.976 $\pm$.003
model	letter			loan		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.646 $\pm$.025	.800 $\pm$.042	.651 $\pm$.048	.921 $\pm$.002	.829 $\pm$.048	.926 $\pm$.008
kNNI	.680 $\pm$.025	.720 $\pm$.053	.911 $\pm$.014	.923 $\pm$.001	.867 $\pm$.047	.947 $\pm$.010
EM	.687 $\pm$.022	.800 $\pm$.042	.951 $\pm$.004	.893 $\pm$.024	.847 $\pm$.051	.927 $\pm$.007
SoftImpute	.652 $\pm$.025	.780 $\pm$.049	.764 $\pm$.028	.922 $\pm$.001	.847 $\pm$.049	.933 $\pm$.007
MICE	.668 $\pm$.024	.850 $\pm$.037	.913 $\pm$.007	.922 $\pm$.002	.852 $\pm$.037	.944 $\pm$.009
missForest	.663 $\pm$.024	.780 $\pm$.049	.806 $\pm$.026	.923 $\pm$.001	.770 $\pm$.045	.919 $\pm$.013
Sinkhorn	.670 $\pm$.024	.720 $\pm$.053	.819 $\pm$.021	.922 $\pm$.002	.783 $\pm$.047	.938 $\pm$.010
GAIN	.680 $\pm$.022	.740 $\pm$.047	.764 $\pm$.030	.919 $\pm$.002	.660 $\pm$.064	.922 $\pm$.017
VAEAC	.727 $\pm$.002	.800 $\pm$.000	.891 $\pm$.002	.828 $\pm$.002	.627 $\pm$.046	.850 $\pm$.015
MIWAE	.641 $\pm$.025	.890 $\pm$.028	.802 $\pm$.019	.920 $\pm$.001	.835 $\pm$.046	.890 $\pm$.010
not-MIWAE	.657 $\pm$.025	.780 $\pm$.051	.781 $\pm$.026	.922 $\pm$.002	.807 $\pm$.052	.937 $\pm$.009
MIRACLE	.691 $\pm$.024	.820 $\pm$.036	.941 $\pm$.008	.923 $\pm$.001	.820 $\pm$.043	.939 $\pm$.010
ReMasker	.677 $\pm$.025	.760 $\pm$.050	.931 $\pm$.010	.923 $\pm$.001	.787 $\pm$.037	.930 $\pm$.013
DrIM _{BASE}	.687 $\pm$.020	.760 $\pm$.050	.937 $\pm$.006	.919 $\pm$.002	.776 $\pm$.048	.898 $\pm$.010
DrIM _{FINE}	.719 $\pm$.016	.892 $\pm$.031	.942 $\pm$.008	.920 $\pm$.002	.787 $\pm$.047	.907 $\pm$.009
model	redwine			whitewine		
	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
Mean	.466 $\pm$.018	.130 $\pm$.127	.787 $\pm$.036	.416 $\pm$.018	.290 $\pm$.097	.485 $\pm$.080
kNNI	.469 $\pm$.020	.145 $\pm$.114	.843 $\pm$.024	.423 $\pm$.017	.295 $\pm$.090	.730 $\pm$.056
EM	.469 $\pm$.021	.083 $\pm$.126	.859 $\pm$.025	.420 $\pm$.019	.320 $\pm$.083	.655 $\pm$.086
SoftImpute	.467 $\pm$.020	.060 $\pm$.149	.874 $\pm$.023	.416 $\pm$.018	.330 $\pm$.082	.811 $\pm$.040
MICE	.470 $\pm$.020	.104 $\pm$.137	.858 $\pm$.023	.417 $\pm$.019	.270 $\pm$.068	.777 $\pm$.049
missForest	.465 $\pm$.021	.129 $\pm$.135	.879 $\pm$.022	.417 $\pm$.018	.280 $\pm$.083	.770 $\pm$.038
Sinkhorn	.468 $\pm$.020	.170 $\pm$.110	.838 $\pm$.020	.419 $\pm$.019	.300 $\pm$.095	.639 $\pm$.103
GAIN	.529 $\pm$.009	.458 $\pm$.088	.895 $\pm$.012	.460 $\pm$.013	.710 $\pm$.048	.800 $\pm$.040
VAEAC	.490 $\pm$.005	.285 $\pm$.062	.804 $\pm$.009	.435 $\pm$.004	.510 $\pm$.043	.803 $\pm$.013
MIWAE	.463 $\pm$.021	.005 $\pm$.139	.813 $\pm$.037	.420 $\pm$.018	.305 $\pm$.081	.741 $\pm$.067
not-MIWAE	.463 $\pm$.019	.062 $\pm$.125	.875 $\pm$.017	.419 $\pm$.018	.330 $\pm$.090	.770 $\pm$.067
MIRACLE	.467 $\pm$.023	.086 $\pm$.151	.864 $\pm$.030	.420 $\pm$.019	.340 $\pm$.093	.671 $\pm$.087
ReMasker	.449 $\pm$.023	.039 $\pm$.138	.745 $\pm$.033	.405 $\pm$.019	.117 $\pm$.148	.572 $\pm$.054
DrIM _{BASE}	.542 $\pm$.007	.503 $\pm$.064	.886 $\pm$.021	.462 $\pm$.010	.870 $\pm$.047	.872 $\pm$.024
DrIM _{FINE}	.564 $\pm$.005	.411 $\pm$.082	.877 $\pm$.012	.491 $\pm$.007	.577 $\pm$.066	.825 $\pm$.034

A.9.4 EFFECT OF CONTRASTIVE LEARNING

Table 17: Effect of contrastive learning under **MCAR**. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. Values in parentheses indicate the performance difference compared to $\text{DrIM}_{\text{BASE}}$, and the **red** highlights the positive improvement.

Rate	model	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
0.2	$\text{DrIM}_{\text{BASE}}$	$.665 \pm .025$	$.692 \pm .026$	$.949 \pm .004$
	$\text{DrIM}_{\text{FINE}}$	$.678 \pm .025$ (+1.95%)	$.690 \pm .029$ (-0.29%)	$.948 \pm .004$ (-0.11%)
0.4	$\text{DrIM}_{\text{BASE}}$	$.602 \pm .024$	$.301 \pm .048$	$.853 \pm .012$
	$\text{DrIM}_{\text{FINE}}$	$.632 \pm .025$ (+4.98%)	$.503 \pm .038$ (+67.11%)	$.889 \pm .011$ (+4.29%)
0.6	$\text{DrIM}_{\text{BASE}}$	$.539 \pm .024$	$.104 \pm .051$	$.708 \pm .024$
	$\text{DrIM}_{\text{FINE}}$	$.588 \pm .025$ (+8.85%)	$.349 \pm .044$ (+233.65%)	$.748 \pm .019$ (+5.64%)
0.8	$\text{DrIM}_{\text{BASE}}$	$.471 \pm .023$	$.027 \pm .050$	$.272 \pm .043$
	$\text{DrIM}_{\text{FINE}}$	$.503 \pm .025$ (+6.81%)	$.176 \pm .055$ (+651.85%)	$.309 \pm .039$ (+13.59%)

Table 18: Effect of contrastive learning under **MNARL**. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. Values in parentheses indicate the performance difference compared to $\text{DrIM}_{\text{BASE}}$, and the **red** highlights the positive improvement.

Rate	model	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
0.2	$\text{DrIM}_{\text{BASE}}$	$.660 \pm .024$	$.658 \pm .031$	$.925 \pm .007$
	$\text{DrIM}_{\text{FINE}}$	$.673 \pm .025$ (+1.97%)	$.711 \pm .026$ (+8.05%)	$.938 \pm .006$ (+1.42%)
0.4	$\text{DrIM}_{\text{BASE}}$	$.591 \pm .023$	$.326 \pm .048$	$.832 \pm .014$
	$\text{DrIM}_{\text{FINE}}$	$.625 \pm .024$ (+5.75%)	$.481 \pm .042$ (+47.55%)	$.864 \pm .010$ (+3.85%)
0.6	$\text{DrIM}_{\text{BASE}}$	$.522 \pm .022$	$.102 \pm .052$	$.652 \pm .029$
	$\text{DrIM}_{\text{FINE}}$	$.573 \pm .024$ (+9.78%)	$.313 \pm .045$ (+205.88%)	$.730 \pm .021$ (+12.00%)
0.8	$\text{DrIM}_{\text{BASE}}$	$.449 \pm .022$	$.049 \pm .051$	$.286 \pm .043$
	$\text{DrIM}_{\text{FINE}}$	$.492 \pm .024$ (+9.57%)	$.235 \pm .054$ (+379.59%)	$.287 \pm .042$ (+0.35%)

Table 19: Effect of contrastive learning under **MNARQ**. The means and the standard errors of the mean across 10 datasets and 10 repeated experiments are reported. $\uparrow$ denotes higher is better. Values in parentheses indicate the performance difference compared to $\text{DrIM}_{\text{BASE}}$, and the **red** highlights the positive improvement.

Rate	model	$F_1 \uparrow$	Model $\uparrow$	Feature $\uparrow$
0.2	$\text{DrIM}_{\text{BASE}}$	$.682 \pm .025$	$.771 \pm .026$	$.945 \pm .005$
	$\text{DrIM}_{\text{FINE}}$	$.689 \pm .025$ (+1.03%)	$.784 \pm .021$ (+1.68%)	$.948 \pm .004$ (+0.32%)
0.4	$\text{DrIM}_{\text{BASE}}$	$.629 \pm .025$	$.534 \pm .037$	$.850 \pm .012$
	$\text{DrIM}_{\text{FINE}}$	$.655 \pm .025$ (+4.14%)	$.618 \pm .033$ (+15.36%)	$.853 \pm .012$ (+0.35%)
0.6	$\text{DrIM}_{\text{BASE}}$	$.583 \pm .026$	$.266 \pm .050$	$.683 \pm .028$
	$\text{DrIM}_{\text{FINE}}$	$.614 \pm .026$ (+5.32%)	$.434 \pm .042$ (+163.27%)	$.748 \pm .021$ (+9.96%)
0.8	$\text{DrIM}_{\text{BASE}}$	$.515 \pm .026$	$.151 \pm .060$	$.545 \pm .033$
	$\text{DrIM}_{\text{FINE}}$	$.572 \pm .025$ (+11.07%)	$.281 \pm .052$ (+86.75%)	$.616 \pm .031$ (+13.19%)

A.9.5 SENSITIVITY ANALYSIS FOR MISSINGNESS RATES

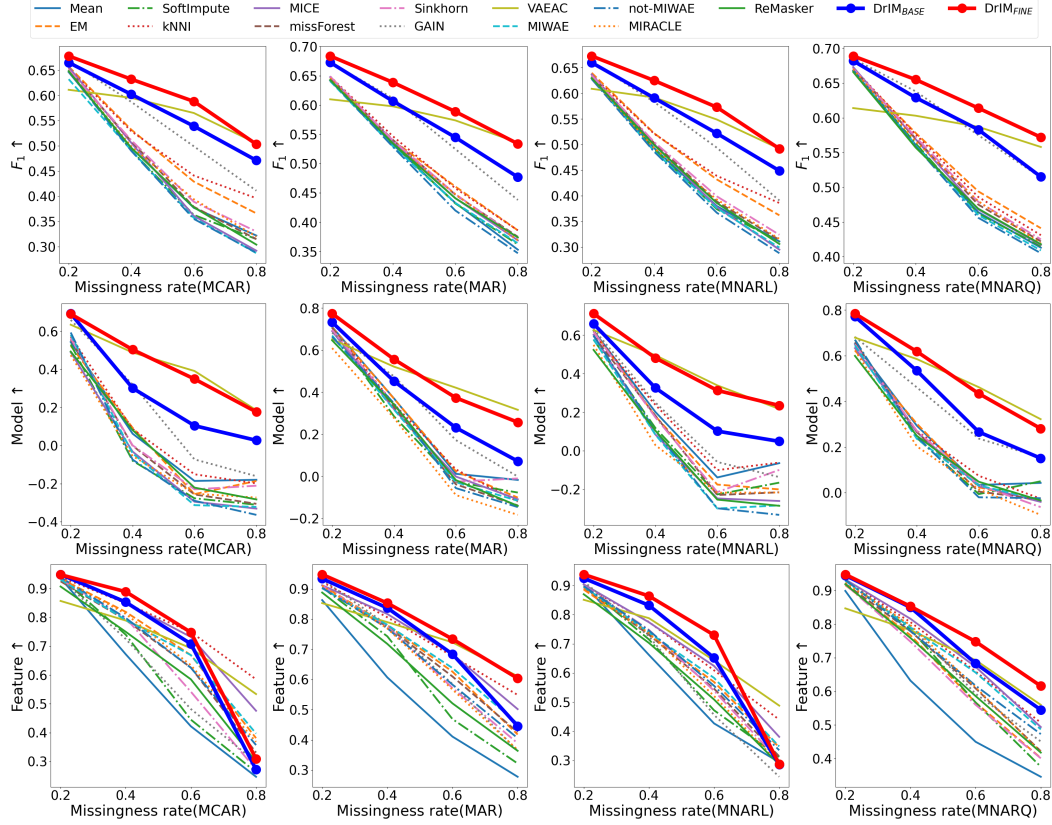


Figure 8: **Sensitivity analysis for missingness rates.** The results MCAR, MAR, MNARL, and MNARQ are shown, with the first row representing classification performance (F_1 score), and the last two rows displaying model selection and feature selection performance. The means of the average across 10 datasets and 10 repeated experiments are reported.