

Process Mining Embeddings: Learning Vector Representations for Petri Nets

Juan G. Colonna¹, Ahmed A. Fares^{2,3}, Márcio Duarte³, and Ricardo Sousa³

¹ Computing Institute (IComp), Federal University of Amazonas (UFAM), Brazil
 juancolonna@icomp.ufam.edu.br

² Faculty of Engineering - University of Porto, Portugal

³ INESC Tec, Porto, Portugal
 {ahmed.a.fares, marcio.s.duarte, ricardo.t.sousa}@inesctec.pt

Abstract. Process mining offers powerful techniques for discovering, analyzing, and enhancing real-world business processes. In this context, Petri nets provide an expressive means of modeling process behavior. However, directly analyzing and comparing intricate Petri net presents challenges. This study introduces PetriNet2Vec, a novel unsupervised methodology based on Natural Language Processing concepts inspired by Doc2Vec and designed to facilitate the effective comparison, clustering, and classification of process models represented as embedding vectors. These embedding vectors allow us to quantify similarities and relationships between different process models. Our methodology was experimentally validated using the PDC Dataset, featuring 96 diverse Petri net models. We performed cluster analysis, created UMAP visualizations, and trained a decision tree to provide compelling evidence for the capability of PetriNet2Vec to discern meaningful patterns and relationships among process models and their constituent tasks. Through a series of experiments, we demonstrated that PetriNet2Vec was capable of learning the structure of Petri nets, as well as the main properties used to simulate the process models of our dataset. Furthermore, our results showcase the utility of the learned embeddings in two crucial downstream tasks within process mining enhancement: process classification and process retrieval.

Keywords: Process Mining · Model enhancement · Embedding vectors
 · Petri Nets

1 Introduction

Modern business process models exhibit a level of complexity that renders traditional analysis tools insufficient for comprehensive understanding and optimization [20]. This challenge demands innovative approaches capable of extracting complex relationships governing process behavior.

Process discovery techniques leverage event data to extract, analyze, and visualize the actual execution of business processes. Generally, a data-driven approach with advanced algorithms is used to construct models, such as Petri

nets, that accurately capture sequences of activities, *i.e.* transitions (activity), and places (state) [16]. Petri nets are mathematical and visual modeling tools used to describe distributed systems. These networks allow us to assess process performance, identify bottlenecks, and uncover hidden patterns through conformance-checking techniques. Ultimately, facilitating process enhancement and optimization.

Traditional process mining core capabilities often fall short when faced with the scale and complexity of modern processes [7]. They may struggle to handle the vast amounts of data generated and lack the sophistication to uncover complex patterns that machine learning techniques could reveal [3].

Embedding vectors are numerical representations of objects or concepts in a continuous vector space [9]. Commonly used in Natural Language Processing (NLP) related tasks, which often involve machine learning algorithms. Embedding vectors capture semantic relationships between entities. In process mining, embeddings can provide a powerful means of representing complex process structures and relationships within process models [19]. For example, we can embed individual activities, control flow structures, or entire process models in vector representations. This enables sophisticated similarity analysis, identification of analogous patterns, and application of predictive modeling techniques that would be difficult to achieve with traditional process representations.

Our proposed approach enables us to encode both the structural information of Process models, in Petri net format, and the individual tasks into compact vector representations (*i.e.* embedding vectors), facilitating various downstream tasks such as similarity analysis and process discovery.

1.1 Problem statement

Given a set $\mathbb{M} = \{M_1, M_2, \dots, M_n\}$ of n Process Models, where each process is represented by a Petri net in PNML format, our goal is to learn a d -dimensional embedding vector $\mathbf{x} \in \mathbb{R}^d$ for each model $M_j \in \mathbb{M}$. The matrix representation of all models will be denoted as $X \in \mathbb{R}^{n \times d}$, with each row representing an embedding vector. Our objective is to capture the structural dependencies of sequential tasks within the processes, facilitating similarity comparisons between models using cosine distance on vector embedding space.

Furthermore, since each model in $\mathbb{M}$ comprises several sequences of tasks and the total number of unique tasks across all processes is denoted as $\mathbb{T} = \{t_1, t_2, \dots, t_k\}$, we also aim to learn an embedding vector $\mathbf{t}_i \in \mathbb{R}^d$ for each task. The matrix of all task embeddings will be denoted as $T \in \mathbb{R}^{k \times d}$, with each row representing an embedding vector. These task embedding vectors serve to capture the inherent characteristics of each task within the process sequences.

2 Related works

Process comparison is crucial in conformance analysis, process enhancement, knowledge transfer, and process retrieval. Existing techniques for process comparison broadly fall into three categories: behavioral analysis, structural analysis,

and task comparison [22]. Behavioral strategies focus on the order of activities within execution logs, while structural approaches analyze the process model as a graph. Task comparison offers a granular view, looking at relationships between individual activities [4].

Existing research on process model similarity analysis offers valuable tools, but certain limitations demand further attention. Behavioral methods can lack effectiveness when variations in terminology mask similarities between processes [6]. This issue is particularly pronounced in localized comparison methods, where direct node-to-node comparisons fail to grasp similarities between models from different domains. Additionally, both behavioral and structural techniques frequently encounter exponential complexity stemming from concurrency and loops within process models [6]. Furthermore, structural approaches reliant on graph edit distance face computational scalability challenges with large graphs due to their NP-Hard complexity [28].

Behavioral process comparison focuses on the execution sequences observed in event logs. Van Der Aalst *et al.* [24] emphasize the value of direct log analysis in uncovering process behavior. Dijkman *et al.* [4] introduce the concept of “causal footprints”, which capture the temporal dependencies between activities to represent process behavior. Kunze *et al.* [10] highlight relationships between exclusivity and strict order as key behavioral aspects. Van Dongen *et al.* [25] develop behavioral metrics for quantifying process similarities and differences. Sánchez-Charles *et al.* [23] extend this by proposing a behavioral distance measure to compare process behavior holistically.

Alternatively, structural methods examine the overall process model as a graph, analyzing elements like node relationships, paths, connectivity, and control flow. Standard techniques utilize edit graph distance to calculate the similarity between process graphs [5]. More nuanced methods consider richer information, such as node types (e.g., AND/XOR) and sequential information [15]. Zhou *et al.* [29] enhance structural techniques by incorporating insights from execution logs and weighting graph edges based on their frequency. While valuable for analyzing overall process flow, such techniques may not fully capture fine-grained behavioral variations that arise from the diverse execution of process activities.

Task comparison provides a detailed analysis of activity relationships within process models. Node-by-node techniques excel at identifying subtle differences, while block-based methods reveal more significant structural changes [12,4,1,27]. Recent clustering approaches hold the potential for comparing localized process sections [18]. However, these localized clustering methods may struggle when analyzing processes with highly variable flows or behavioral patterns.

In the domain of Similarity-based retrieval of semantic graphs, Hoffman *et al.* [8] proposed an approach employing Graph Neural Networks (GNNs) for process mining, framing it as a graph matching challenge. Despite the inherent complexity of GNNs, necessitating larger datasets for effective training, their methodology shows significant promise, particularly in retrieval scenarios essential to Process-Oriented Case-Based Reasoning (POCBR) within smart manu-

facturing - a focus distinct from our current study. Nonetheless, in Section 5.4, we comprehensively evaluate our methodology, specifically in the realm of semantic retrieval tasks for process models, utilizing *PetriNet2Vec*. This assessment underscores the resilience and adaptability of our retrieval approach, highlighting its potential for broader applications in the field.

The optimal comparison strategy hinges on the desired analytical depth and specific research questions. While current techniques offer valuable insights, their limitations suggest the need for a method that offers a more holistic view of process execution, considering both behavioral patterns and the broader process structure. Our proposed approach addresses these needs. Firstly, it considers the entire process and its structure, mirroring structural approaches. Secondly, it enables the comparison of multiple processes and clusters, echoing the strengths of most behavioral approaches.

3 Background

3.1 Learning embeddings with *doc2vec* and *graph2vec*

The *doc2vec* methodology [11], an extension of the *word2vec* concept, has become popular for enabling the simultaneous learning of embedding vectors for both documents and individual words. In *word2vec* [14], a word is predicted based on the words in its context. Figure 1 illustrates the functioning of the CBOW (Continuous Bag-of-Words) approach. While this figure is useful for illustrating the *word2vec* concept, we can reinterpret the optimization objective function during training as a binary classifier, predicting $P(1|w_i, w_c)$ when the word w_i is part of the word context $w_c = \{\dots, w_{i-2}, w_{i-1}, w_{i+1}, w_{i+2}, \dots\}$. Thus, if w_i is not in the context w_c , then we have $P(0|w_i, w_c)$. Then, during training, the *negative sampling* strategy generates positive and negative example pairs (w_i, w_c) .

Extending this formulation to learn embeddings for documents is straightforward. We can add a vector d_n to the context w_c , treating it as a new word representing the entire document (Figure 1b). Thus, the optimization objective that allows learning an embedding for this ‘new word’ becomes predicting $P(1|w_i, w_c, d_n)$ when the word w_i belongs to the context w_c and has also been sampled from the document d_n , otherwise $P(0|w_i, w_c, d_n)$ when w_i is not present in the document d_n .

The *graph2vec* algorithm is a powerful formulation that allows us to learn to embed vectors for graph representation [17]. Given a graph $G_n = (V, N)$ consisting of a set of vertices V and edges N , *graph2vec* naturally extends the principles of *doc2vec*. By analogy, each node in the graph can be seen as a word $V_i \equiv w_i$, with the neighboring nodes $V_c \equiv w_c$ directly connected to V_i , serving as its context. Consequently, if each node is associated with a token t_k , as depicted in Figure 1c, we can concurrently learn embeddings for the entire graph h_n and for each individual node.

This approach outperforms other graph embedding techniques because it trains on various graph samples and incorporates neighborhood information for

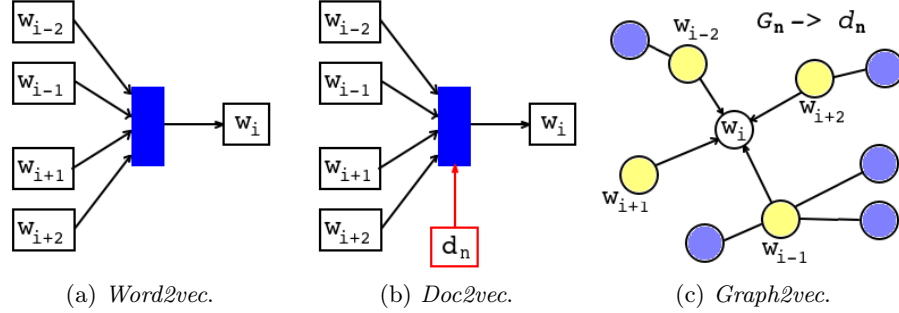


Fig. 1. Subfigure 1a depicts the CBOW approach for *word2vec*. Subfigure 1b demonstrates the incorporation of the document ID from which the words were sampled. Subfigure 1c illustrates a graph where the yellow nodes represent the ‘context’ of node w_i , while adhering to the same nomenclature used in *doc2vec* for the words within the context of w_i .

every node. Thus, it helps the neural network learn complex graph structures more effectively. As a result, *graph2vec* embeddings can accurately capture similarities between graphs with similar structures, making them highly useful across different downstream applications [2].

3.2 Cluster algorithm

As discussed in the previous sections, one of the objectives of this work is to learn embedding for each Petri net model and group them by similarity. For this purpose, we adopted a version of the Hierarchical Density-Based Spatial Clustering of Applications with Noise (HDBSCAN) clustering algorithm with cosine distance [13].

HDBSCAN stands out as a sophisticated clustering algorithm capable of effectively identifying clusters of various shapes and densities. Its hierarchical clustering methodology allows it to autonomously discern clusters at various density levels without prior knowledge of the ideal cluster count, providing flexibility in estimating the number of clusters present. Furthermore, thanks to implementing cosine distance, HDBSCAN demonstrates proficiency in handling high-dimensional data spaces, such as those found in text data or feature vectors.

Cosine similarity is a robust measure for clustering sparse data commonly encountered in vector embedding, accommodating fluctuations in vector magnitudes. This characteristic makes it particularly suitable for tasks where traditional techniques like k-means fail, especially when dealing with non-Euclidean distances such as cosine similarity. Unlike k-means, which faces challenges in combining the average centroid calculation with cosine distance due to their inherent incompatibility, HDBSCAN with cosine distance presents a versatile and effective solution for clustering sparse, high-dimensional data, positioning it as the chosen method for our application.

The output of HDBSCAN can be combined with the Silhouette score to evaluate the quality of formed clusters when ground truth labels are unavailable. This coefficient measures how well each data point fits into its assigned cluster, with values ranging from -1 to 1. A score near 1 indicates that data points are well within their clusters and far from others, while negative values close to -1 suggest poor cluster fit. Scores around 0 imply potential cluster overlap. Additionally, the average Silhouette score serves as a global measure, providing insights into the overall quality of the formed clusters.

3.3 Process models dataset

Our research leverages the PDC Dataset [26], specifically curated for the Process Discovery Contest. This dataset encompasses a total of 96 Petri net models stored in PNML format. As a specialized class of models, Petri nets present a versatile range of configurations, rendering them indispensable for evaluating process discovery algorithms and techniques.

The dataset is derived from a base model named `pdc2023_000000.pnml` and spans various configurations denoted by specific letters from A to F, such as `pdc2023_ABCDEF.pnml`. Each letter stands for a different configuration parameter:

- A: Dependent Tasks (Long-term Dependencies), configured as 0 (No) or 1 (Yes). If Yes, bypass connections are added to shortcut long-term dependent transitions;
- B: Loops, configured as 0 (No), 1 (Simple), or 2 (Complex), determining the treatment of transitions initiating loops and shortcuts between the loop and main branches;
- C: OR constructs, configured as 0 (No) or 1 (Yes), controlling transitions involving only inputs or generating only outputs for OR-joins and OR-splits;
- D: Routing Constructs (Invisible Tasks), configured as 0 (No) or 1 (Yes), making certain transitions invisible when set to Yes;
- E: Optional Tasks, configured as 0 (No) or 1 (Yes), enabling the skipping of visible transitions with the addition of invisible transitions when set to Yes; and
- F: Duplicate Tasks (Recurrent Activities), configured as 0 (No) or 1 (Yes), involving the relabeling of transitions to account for duplicate tasks when set to Yes.

All candidate models were systematically generated following a sequential rule that applies each parameter to the base model. In essence, the characters in this naming convention serve as flags that, when activated, apply these specific rules to the process model. For instance, the first four candidate models are named `pdc2023_000000.pnml`, `pdc2023_000001.pnml`, `pdc2023_000010.pnml`, and `pdc2023_000011.pnml`. This nomenclature signifies that the second, third, and fourth models are variations of the base model, incorporating Duplicate Tasks (F), Optional Tasks (E), or both (F and E), respectively. Subsequently, the next configuration D (`pdc2023_000100.pnml`) is activated, and the procedure repeats.

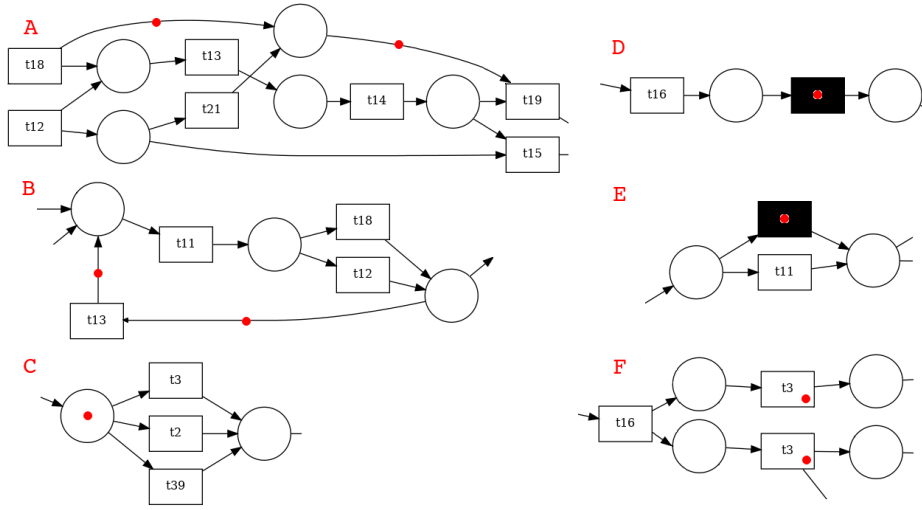


Fig. 2. Illustration of the rules applied to a process model. Red dots indicate the effect caused by each rule. (A) shows a bypass connection, (B) shows a loop, (C) shows an OR-construct, (D) shows an invisible task, (E) shows an optional task, and (F) shows a duplicated task and demonstrates an AND split.

Figure 2 illustrates through Petri net segments where these rules have been applied. This sequential generation process encompasses all possible combinations, where the 2^5 binary combinations are multiplied by the three configurations of rule B, resulting in a total of $2^5 * 3 = 96$ model variations. Understanding this generation procedure is fundamental to grasping how our learning algorithm identifies clusters with similar models generated by comparable combinations of these rules.

4 Methodology for Learning Petri Net Embeddings

Consider the Petri net illustrated on the left side of Figure 3 as an example of a process model. In this representation, each box represents a transition, and the circles represent states. The black dot is a token that traverses this network from left to right, visiting all states along a path. A transition must be completed to progress from one state to the next. Therefore, transitions represent tasks that need to be accomplished. Transitions are pivotal elements in these representations as they denote the activities that constitute the entire process. Transitions can be labeled using natural language. Here, for consistency and simplicity, we have labeled them using IDs, namely as t_1 , t_2 , and so forth.

In a Petri net, paths are represented by directed arrows, which also denote temporal dependencies. For instance, the transition t_3 cannot be initiated until the transition t_1 is completed. However, some paths can be executed in parallel; for example, transition t_3 can be executed concurrently with t_2 . This indicates

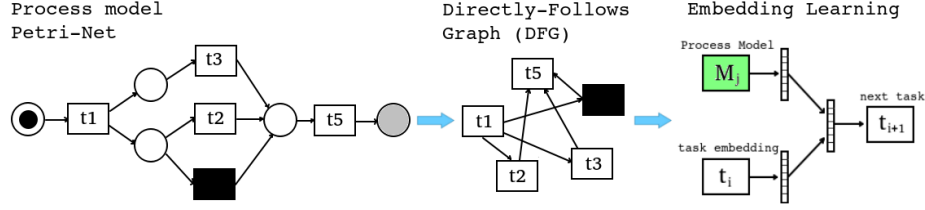


Fig. 3. Proposed methodology. On the left, a Petri net representation of a process model; in the middle, an equivalent representation as a directed graph of transitions; on the right, a Distributed Memory algorithm used for jointly learning embeddings for the process and tasks.

that a predecessor transition, like t_1 , acts like an AND split, where both resultant paths are taken. Conversely, if a state splits into two paths, it is interpreted as an XOR split, meaning that only one of the two paths is chosen, such as t_2 or the black transition, but never both. Finally, the black box is a *Silent* transition and functioning as a wildcard. In this example, the black box enables skip connections in the model *i.e.* permitting optionally execution of t_2 .

Our hypothesis is that since every transition can be represented in natural language, and there is a mandatory order dependency between these transitions, then a path inside the model can be considered equivalent to a sentence in a text document. Therefore, the set of all possible paths inside the model is equivalent to a text paragraph. Hence, without loss of generality, we can assume that the *doc2vec* methodology can be employed here to learn a vector embedding for each process model in our dataset. Indeed, we can also jointly learn a vector embedding for every task by considering the set of all tasks in all models.

The embedding vectors generated by *doc2vec* carry semantic meanings about documents and words. Here, these vectors will convey information and meaning about the process model and about every task across all models. Thus, those vectors may be useful for downstream tasks, such as model comparison, clustering, or classification. Therefore, by analogy, we named our methodology as *PetriNet2Vec*.

Despite its simplicity, a key step in our methodology is mapping a process model into an equivalent representation suitable for algorithm training and embedding vector learning. Here, we introduce the concept of *graph2vec*. Prior to training, we generate a document list containing tasks. To achieve this, we transform each Petri net into an intermediate representation resembling a Directly-Follows Graph (DFG). DFGs offer a popular and comprehensible format for representing business processes in process mining. Figure 3 (center) illustrates a DFG derived from the Petri net on the left. In this graph, we remove places and directly connect tasks, forming a classic directed graph where each node represents a task. Using that DFG, we construct tuples (t_i, t_{i+1}, m_j) , where m_j is a unique identifier for the j -th process model in the dataset.

To train *PetriNet2Vec*, we represent a model as the set of all possible transition pairs within it. Furthermore, we adopted the Distributed Memory (DM) algorithm implemented by the Gensim [21] library in Python, as illustrated on the right side of Figure 3. The diversity of process models in the training dataset ensures that some pairs of transitions are more frequent in some process models than others, and some transition pairs may even exist in some models but not in others. Therefore, by training using the described tuples, we can ensure that the learned vectors represent every model uniquely, carrying information about their structure. During training, we aim to maximize the probability of task t_{i+1} being in the context of task t_i and model m_j . We train on the entire dataset, maximizing the average probability loss:

$$\text{maximize: } \frac{1}{M} \sum_j \frac{1}{T} \sum_i \log P(t_{i+1}|t_i, m_j), \quad (1)$$

where T represents the set of all tasks and M denotes the set of all models involved.

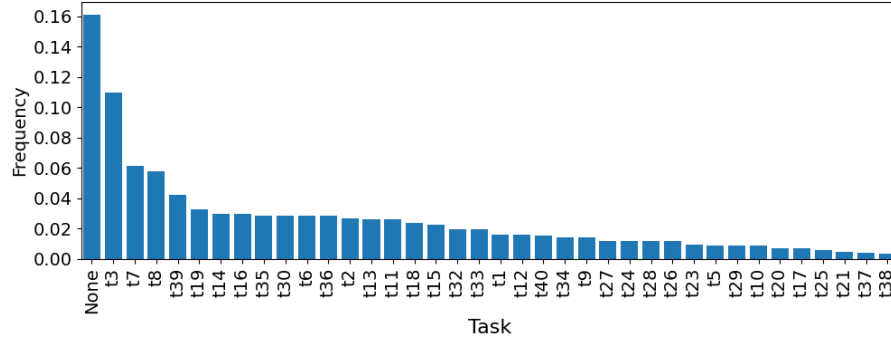
This training method does not require any model labels or supervision; in other words, this embedding training is unsupervised. To improve the separability between the learned vectors, we adopted the negative sampling strategy [14]. With this approach, in every training epoch, some pairs are randomly sampled from other models to form negative tuples for training that do not exist in the current model being trained. With negative sampling, the maximization objective becomes $P(1|t_{i+1}, t_i, m_j)$, *i.e.*, predict 1 when the task t_{i+1} is in the context of task t_i and model m_j . Nevertheless, all these sampling and training epochs occur in parallel for speed and efficiency improvements.

5 Results

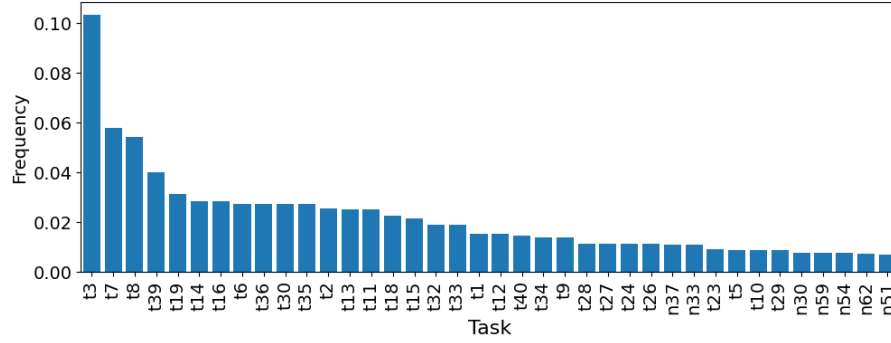
5.1 Cluster analysis

As previously discussed, specific tasks within our Petri net models may be unspecified, denoted by black boxes. We face two choices in developing our methodology: we can adopt a generic and uniform name for all black boxes in all models, for instance, using the ‘None’ token, akin to an out-of-vocabulary token in *word2vec*, or uniquely label each black box according to its position in the network, referred to as ‘ n_x ’ transitions from Figure 4. These alternatives for constructing the task dictionary (also known as the dictionary of tokens) influence the embedding vectors we aim to learn during *PetriNet2Vec* model training. Thus, our initial inquiry delves into how the quality of the clusters is impacted under each scenario.

Figure 4 illustrates the histograms depicting task distributions across all models when employing the ‘None’ token versus individually naming black boxes. A comparative analysis reveals stark differences in task frequencies. Utilizing unique names for black boxes necessitates the Distributed Memory (DM) algorithm to learn a more extensive set of embedding vectors, given the expanded



(a) Methodoly 1 ('None')



(b) Methodoly 2

Fig. 4. Caption for this figure with two images

token dictionary. While this might offer an advantage in distinguishing structures between models, our relatively modest dataset for training implies that the resultant vectors may suffer in quality. In essence, attempting to learn more embedding vectors without sufficient sample Petri net models compromises the final cluster quality across models. Therefore, due to these constraints, methodology 1, which entails a more straightforward approach with a smaller token dictionary, has a better chance of converging to more discriminative embedding vectors. This assertion finds validation in the results presented in Table 1.

Table 1 compares the silhouette scores for the two methodologies. The left-most column denotes the size of embeddings d_n , while the adjacent columns vary the hyperparameter named ‘minimum cluster size’ within the HDBSCAN clustering algorithm. Lower minimum cluster sizes elevate the probability of forming smaller clusters, whereas higher values tend to yield larger, amalgamated clusters. Given the silhouette score’s sensitivity to cluster size, caution is warranted in adjusting this parameter beyond these bounds. After numerous iterations, we identified an acceptable range between 2 and 4.

Silhouette scores were computed by running each methodology ten times, with 1000 training epochs for each investigated embedding size. The central values in the table represent the averages across these runs, while values in parentheses are standard deviations. Consistently, Methodology 1 yields superior silhouette scores compared to Methodology 2, indicating more cohesive clusters. Optimal results were attained with $d_n = 8$ and a minimum cluster size of four. However, slight overlaps in standard deviations across the $d_n = 8$ row suggest that any minimum cluster size, *i.e.*, 2, 3, or 4, could be considered viable for grouping PNML models using HDBSCAN and cosine distance. Henceforth, we will adhere to and explore Methodology 1, employing embeddings of size $d_n = 8$ and a minimum cluster size of 4.

Table 1. Comparison of silhouette scores between Methodology 1 and Methodology 2. Silhouette scores were averaged across ten runs with 1000 training epochs for each embedding size. Methodology 1 consistently yields superior results, particularly evident with an embedding size of $d_n = 8$.

d_n	Methodology 1 (‘None’)			Methodology 2		
	2	3	4	2	3	4
4	0.61 (0.10)	0.69 (0.04)	0.69 (0.04)	0.56 (0.05)	0.62 (0.04)	0.51 (0.07)
8	0.71 (0.02)	0.72 (0.02)	0.73 (0.01)	0.48 (0.03)	0.45 (0.02)	0.25 (0.03)
16	0.47 (0.01)	0.45 (0.03)	0.38 (0.03)	0.32 (0.02)	0.32 (0.02)	0.21 (0.02)
32	0.42 (0.02)	0.41 (0.00)	0.34 (0.02)	0.29 (0.01)	0.28 (0.01)	0.17 (0.01)

5.2 Visual Cluster Assessment

After determining the size of the embedding vectors and running the HDBSCAN algorithm with cosine distance, it was found that the models naturally cluster

into nine groups. The silhouette plot depicted in Figure 5a illustrates the distribution of models across these clusters. Higher silhouette values indicate better inner cluster cohesion. The dashed vertical red line represents the average silhouette across all clusters, measuring the overall quality of the formed groups. Clusters with more members appear wider. The smallest cluster size identified by HDBSCAN was eight, while the largest was seventeen.

Since the process models are represented by vectors of size 8, it is necessary to employ dimensionality reduction techniques to visualize them. We chose the UMAP technique to project these models into two dimensions, as shown in Figure 5b. In this scatter plot, each point represents a process model (a Petri net) from our dataset. The colors of the points correspond to the clusters identified in the Silhouette plot. We opted for UMAP for two reasons: its 2D projection is non-linear and can be performed using cosine distance. This figure demonstrates excellent quality and separation among the groups formed by the embedding vectors. However, it is important to note that we are observing a distorted representation of an 8-dimensional vector space, so some groups of points may not be well accommodated in 2D.

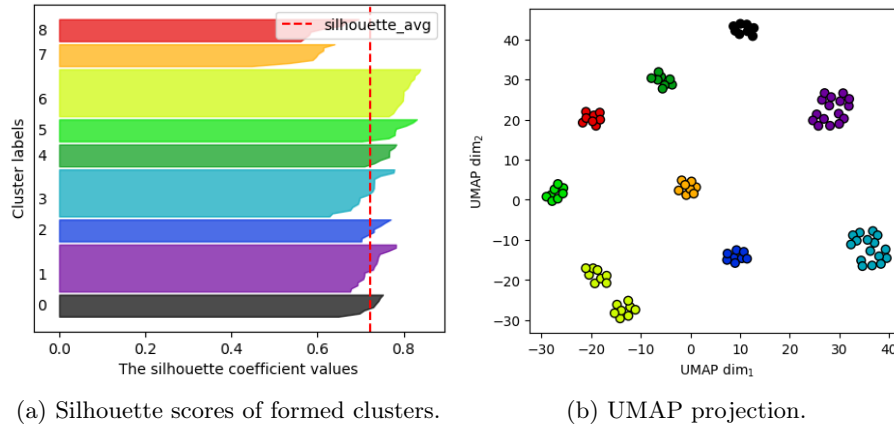


Fig. 5. Visual assessment. Left: Silhouette plot of model clusters with average silhouette indicated. Right: UMAP projection of process models with cluster colors.

Upon closer examination within the cluster members, we observed emerging patterns related to the formation rules used to generate this dataset, as described in Section 3.3. For example, the process models within cluster C_0 include Petri nets⁴: ‘000000’, ‘000010’, ‘001000’, ‘001010’, ‘100000’, ‘100010’, ‘101000’, ‘101010’, all of which lack loops (rule $B = 0$), invisible tasks (rule $D = 0$), and duplicated tasks (rule $F = 0$). We observed similar patterns in other clusters,

⁴ These are abbreviated names of the models; for instance, ‘pdc2023_101010.pnml’ is shortened to ‘101010’.

characterized by different combinations of these rules. This raises an interesting question: Can we identify the formation rules common to all cluster members? To investigate this further, we fitted a decision tree using learning embedding as feature vectors and the cluster identifier as labels. The resulting tree is shown in Figure 6. This tree perfectly fits our clusters, revealing the common formation rules shared by all members of each formed cluster. For instance, all members of cluster C_7 lack duplicated tasks (rule $F = 0$) but have invisible tasks (rule $D = 1$), loops (rule $B = 1$ or $B = 2$), and OR-splits (rule $C = 1$). This strong coherence among cluster members indicates that our methodology discovers the structural properties of the process models.

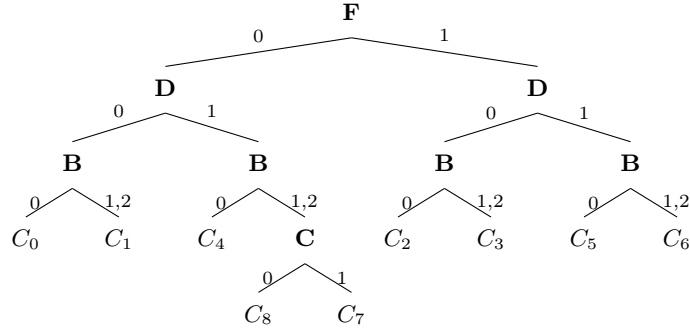
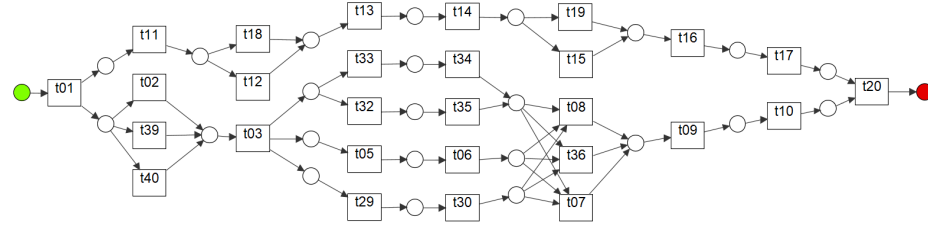


Fig. 6. Decision Tree revealing common formation rules within clusters. The most discriminative rule F , which distinguishes between having or not having duplicated tasks, serves as the root of this tree.

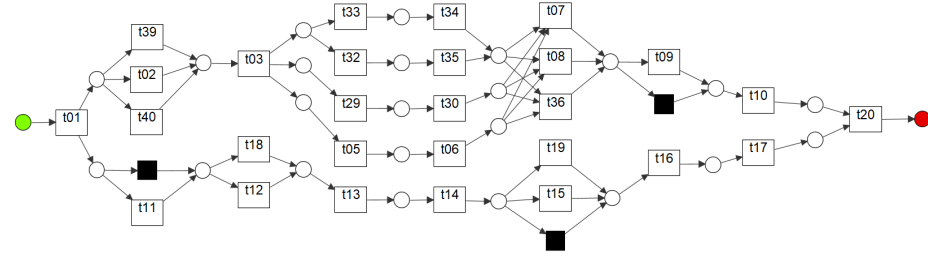
Finally, Figure 7 displays the first four members of cluster C_0 , including the base model and their neighbors. From these models, it can be observed that none of them contain loops, invisible tasks, or duplicated tasks.

5.3 Expanding Our Cluster Analysis to Task Embeddings

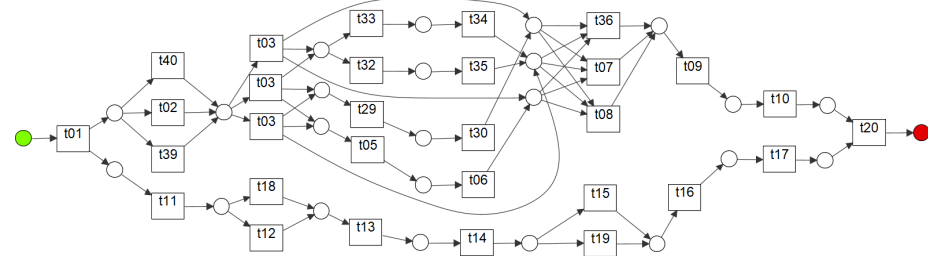
Until now, our analysis has primarily focused on the embeddings of models. However, during *PetriNet2Vec* training, we also learn task embeddings, meaning that each task depicted in Figure 4a is also mapped to an 8-dimensional embedding vector. Employing HDBSCAN on these task embeddings revealed the emergence of five clusters. Figure 8a displays the Silhouette plot, while Figure 8b illustrates the 2-dimensional projection of these task vectors. Unlike the clusters of models, the clusters of tasks exhibit lower quality, and some tasks were not associated with any cluster, indicated as white dots (or cluster noises). This occurred because the cosine similarity with their closest neighbor was considered too large by the HDBSCAN algorithm. Nonetheless, they still exhibit some degree of cohesion, and most importantly, they play a significant role during the training of



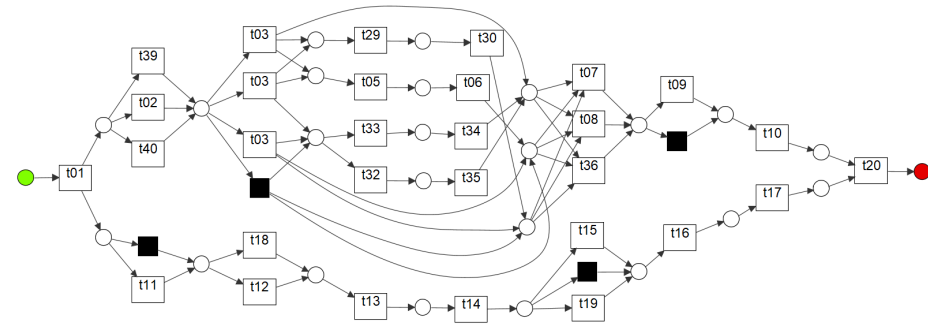
(a) Base model '000000' (no active rules).



(b) '000010' (optional tasks).



(c) '001000' (OR-constructs).



(d) '001010' (optional tasks and OR-constructs).

Fig. 7. Process models members of cluster C_0 .

model embeddings, granting *PetriNet2Vec* greater flexibility in accommodating those model embeddings.

Upon examining the scatter plot in Figure 8a, we observed a tendency for tasks sharing OR-splits, such as t_2 , t_{39} , and t_{40} , to be grouped together, likely owing to their resemblance to text synonyms in *doc2vec*. By analogy, tasks preceding an OR-split create alternate paths within the models, enabling the process to achieve the same subsequent state through diverse routes. Intriguingly, some white tasks are frequently subject to substitution by black boxes when the $D = 1$ rule is applied. As these tasks are replaced, their occurrence diminishes in PNML files, consequently reducing their sampling during training. Moreover, infrequent tasks exist in models that are not necessarily substituted by black boxes. Analogously, these tasks mimic the behavior of less frequent words in *doc2vec*. This becomes apparent when contrasting the histogram in Figure 4a with Figure 8b, where less frequent tasks are predominantly denoted as white dots. Nonetheless, while other analogies may be drawn, it is crucial to acknowledge that certain cluster artifacts may arise due to the inherent limitations of HDBSCAN.

To extend our analysis beyond the clusters produced by HDBSCAN, we also examined the similarity matrix depicted in Figure 9. This matrix represents the cosine similarity between every pair of tasks. The stronger the color, the higher the similarity. From this matrix, we can address questions like: What are the most similar (or dissimilar) tasks to a given task in vector space? Since we use a task dictionary with only task IDs, giving them meaning is challenging. However, in a dataset with real business process models, these IDs are translated into interpretable tasks that must be accomplished to navigate the Petri net process flow. For instance, from this matrix, we can identify the set of tasks most similar to t_{36} are t_7 and t_8 , and the most dissimilar are t_{13} and t_{25} . This highlights that t_{13} and t_{25} lie in a parallel branch to t_{36} .

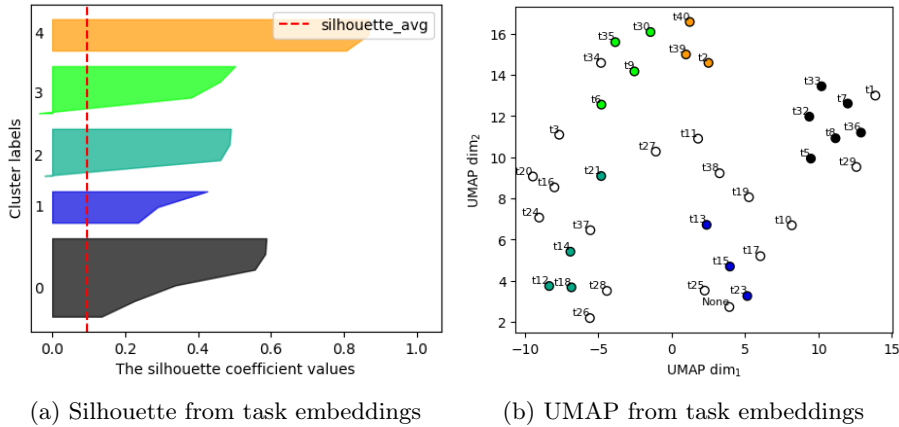


Fig. 8. Visual assessment of task embeddings.

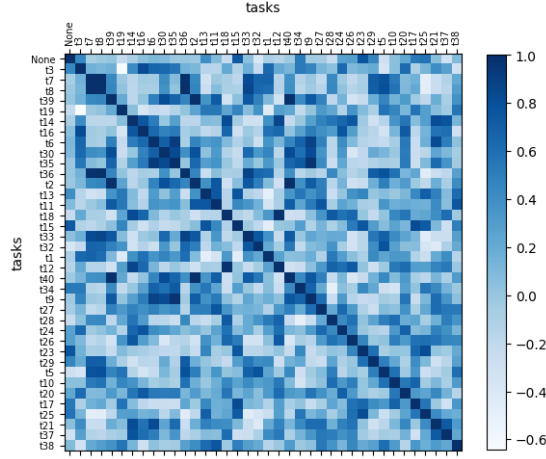


Fig. 9. Cosine similarity matrix illustrating relationships between task embeddings.

5.4 Examples of Downstream Tasks

Once our methodology is complete and the *PetriNet2Vec* algorithm has been trained, we can tackle several downstream tasks using the learned embeddings, such as model or task query for similarity, model or task classification, and more. Also, we can explore what *PetriNet2Vec* algorithm has already learned by examining the relationships between task embeddings and model embeddings.

Our first example of the downstream task is the *model query*, which involves randomly selecting a model, obtaining its embedding vector, and consulting the most similar model in the database by comparing the cosine distance between the query and each model represented by an embedding vector. Figure 10 demonstrates this process retrieval task, where model 20 ('010100') was the query and model 22 ('010110') was the most similar recovered model. Upon comparing these two process models, we observe that the only difference is the activation of rule E, optional paths, in the returned model. This approach can also be adapted to return multiple answers and rank them accordingly.

Addressing the complexity and significance of querying business process models provides invaluable insights for large enterprises managing numerous processes. It enables them to pinpoint structurally similar processes, facilitating informed decisions such as merging or replacing process models. Moreover, this methodology can be extended to identifying similar tasks within processes. Figure 11 visually demonstrates an investigation of this nature, illustrating the similarity between each task embedding and each model embedding. In this figure, the models are ordered according to the clusters they formed during the application of HDBSCAN. Notably, certain task embeddings closely align with all model embeddings within the same cluster, unveiling compelling patterns. Nevertheless, interpreting task-model embeddings demands caution. Although they

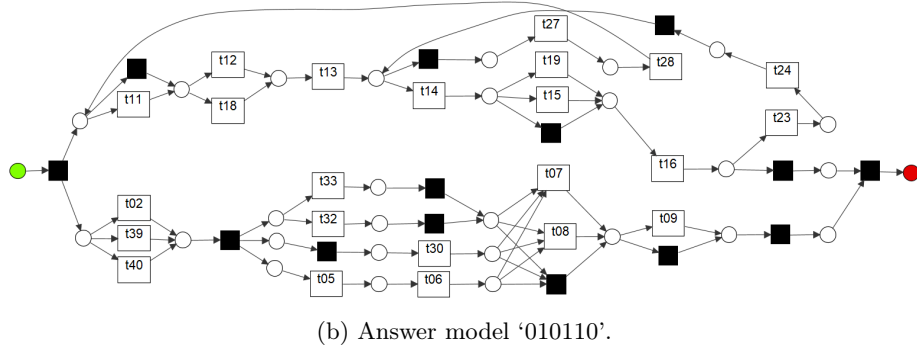
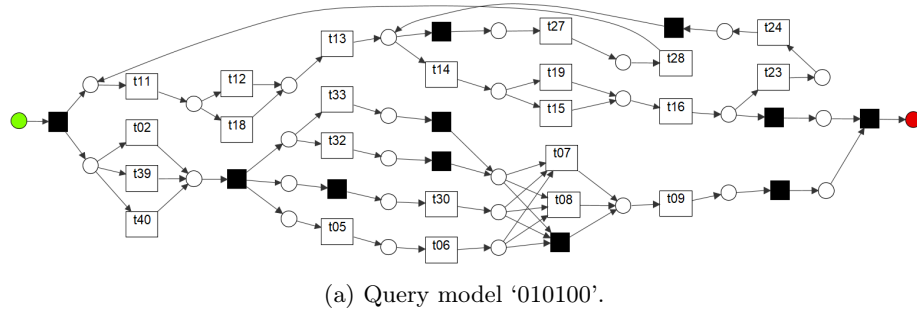


Fig. 10. Example of solving a model query using cosine similarity. In this case, the red dots in (b) highlight the only two differences between the query and the returned models, providing a clear visual representation of the similarities and differences between them.

share the same embedding dimensions, they are different vector spaces. Consequently, the resulting score would not directly translate to a strict interpretation like model-to-model similarity. Nevertheless, it still offers valuable insights into the task’s relevance to the model’s knowledge.

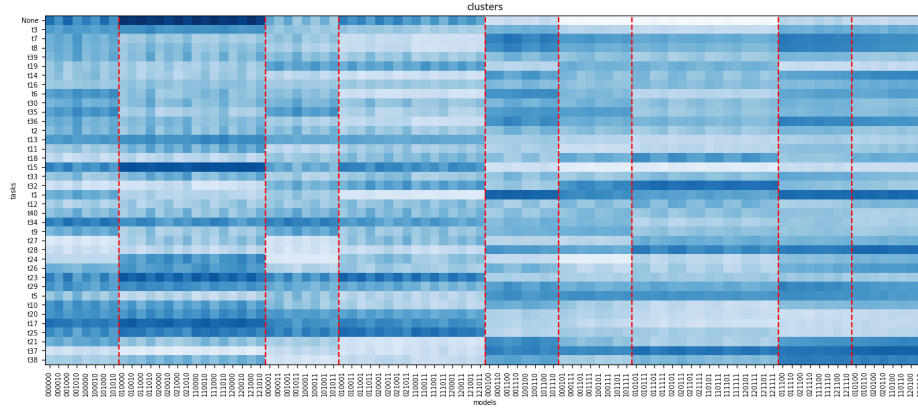


Fig. 11. Task and model embeddings’ cosine similarity visualization. Models are grouped by clusters, separated by vertical dashed lines. This visualization highlights the relationships learned between tasks and models, facilitating the identification of similar processes and tasks within the business process model.

Our second example of a downstream task for model embeddings is classifying processes based on formation rules. In this case, we train a classifier using the embedding vectors as features and the rules described in Section 3.3 as labels. Six k -NN classifiers were trained using 5-fold cross-validation and cosine distance. The results can be found in the confusion matrices shown in Figure 12. In this figure, we can observe that some properties of the models, mainly those that achieved higher accuracy, were effectively incorporated by our methodology when creating the embeddings, such as loops, OR-constructs, invisible tasks, and duplicated tasks. However, bypass connections (representing long-term tasks) and optional tasks were more challenging to recognize from the learned embedding, although the results can still be considered satisfactory. This limitation probably arises from the fact that our methodology only considers pairs of consecutive tasks, neglecting, for example, long-term dependencies.

Finally, these examples demonstrate that, in a real-world scenario, our methodology could be used to identify properties of business process models.

6 Conclusions

In our study, we introduce a novel approach for mapping process models, represented as Petri networks, into embedding vectors. Drawing inspiration from

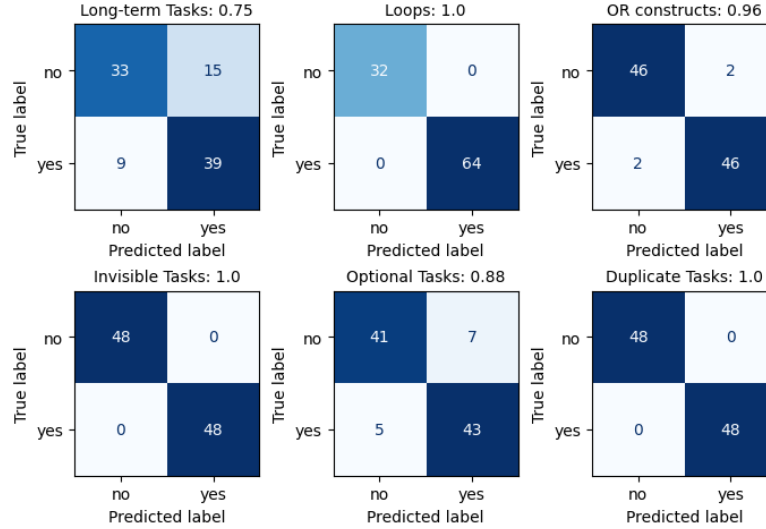


Fig. 12. Confusion matrices of 1-NN with cosine distance trained to recognize the rules A, B, C, D, E, and F from the model embeddings. Each subfigure title indicates the accuracy achieved using the 5-fold cross-validation methodology.

doc2vec, our approach inherits the advantageous characteristics of these methods while offering its own simplicity and effectiveness in the realm of Process Mining. Our hypothesis posits that process paths resemble sentences in text documents, allowing for the application of *doc2vec* methodology. This analogy facilitates the learning of vector embeddings for both process models and individual tasks, leading to the development of the *PetriNet2Vec* methodology.

The training process aims to maximize the probability of task sequences within models. The unsupervised nature of the training, coupled with negative sampling, ensures that robust vector embeddings are learned, capturing the nuances of different process models. Through a series of experiments, we demonstrated that our *PetriNet2Vec* method was capable of learning the structure of Petri nets, as well as the main properties used in constructing the process models in our dataset. Furthermore, our results showcase the utility of the learned embeddings in two crucial downstream tasks within process mining: process classification and process retrieval.

Further analysis within clusters revealed common formation rules learned by *PetriNet2Vec*, as depicted by the Decision Tree, indicating strong coherence among cluster members. For instance, examination of cluster C_0 (Figure 7) showcased models with shared structural properties, such as the absence of loops, invisible tasks, or duplicated tasks. Additionally, exploring task similarity via cosine similarity matrices provided insights into task relationships in vector space, despite the challenge of interpreting tasks solely by their IDs. By identifying tasks with the highest and lowest similarities to a given task, we gained valuable

insights into parallel branches and potential structural relationships within the process flow.

The code used in these experiments can be found in our Github repository for reference⁵. Additionally, our PetriNet2Vec package is now conveniently accessible as a Python package, installable via the **pip** tool⁶. In our future work, we plan to conduct a case study using real-world data collected from a company, as opposed to relying on simulated datasets. This will allow us to validate our findings in a practical, business environment and enhance the applicability of our research. Additionally, we plan to utilize these embeddings to enhance process discovery. For instance, we will compare the effectiveness of our approach with that of a baseline model, aiming to demonstrate improvements in accuracy and efficiency. Furthermore, we plan to expand our methodology by incorporating additional tasks into the model's context. For example, this extension could involve modifying Equation 1 to include $t_{i-1}, t_{i-2}, t_{i-3}, \dots$, thereby integrating a broader task context into the model. This enhancement aims to capture deeper temporal dependencies, potentially improving the model's predictive accuracy and relevance in complex scenarios.

Acknowledgements

This paper is funded/financed in total by Component 5 - Capitalization and Business Innovation, integrated in the Resilience Dimension of the Recovery and Resilience Plan within the scope of the Recovery and Resilience Mechanism (MRR) of the European Union (EU), framed in the Next Generation EU, for the period 2021 - 2026, within project FAIST, with reference 66.

It also supported by National Funds through the Portuguese funding agency, FCT - Fundação para a Ciência e a Tecnologia, within project LA/P/0063/2020 (DOI 10.54499/LA/P/0063/2020) under INESC TEC International Visiting Researcher Programme.

This work was supported by the Coordenação de Aperfeiçoamento de Pessoal de Nível Superior - Brazil (CAPES-PROEX) - Code of Financing 001, and also by the Fundação de Amparo à Pesquisa do Estado do Amazonas - FAPEAM - through the POSGRAD 23-24 project.

References

1. Bae, J., Caverlee, J., Liu, L., Yan, H.: Process mining by measuring process block similarity. In: Business Process Management Workshops: BPM 2006 International Workshops, BPD, BPI, ENEL, GPWW, DPM, semantics4ws, Vienna, Austria, September 4-7, 2006. Proceedings 4. pp. 141–152. Springer (2006)
2. Cai, H., Zheng, V.W., Chang, K.: A comprehensive survey of graph embedding: Problems, techniques, and applications. *IEEE Transactions on Knowledge and Data Engineering* **30**(09), 1616–1637 (sep 2018). <https://doi.org/10.1109/TKDE.2018.2807452>

⁵ <https://github.com/juancolonna/PetriNet2Vec>

⁶ <https://pypi.org/project/PetriNet2Vec>

3. Cárdenas Maita, A.R., Martins, L.C., López Paz, C.R., Rafferty, L., Hung, P.C.K., Peres, S.M., Fantinato, M.: A systematic mapping study of process mining. *Enterprise Information Systems* **12**(5), 505–549 (2018). <https://doi.org/10.1080/17517575.2017.1402371>, <https://doi.org/10.1080/17517575.2017.1402371>, publisher: Taylor & Francis _eprint: <https://doi.org/10.1080/17517575.2017.1402371>
4. Dijkman, R., Dumas, M., Dongen, B.v., Käärik, R., Mendling, J.: Similarity of business process models: Metrics and evaluation. *Information Systems* **36**(2), 498–516 (2011). <https://doi.org/https://doi.org/10.1016/j.is.2010.09.006>, <https://www.sciencedirect.com/science/article/pii/S0306437910001006>
5. Dijkman, R., Dumas, M., García-Bañuelos, L.: Graph matching algorithms for business process model similarity search. In: *Business Process Management: 7th International Conference, BPM 2009, Ulm, Germany, September 8-10, 2009. Proceedings 7*. pp. 48–63. Springer (2009)
6. Dijkman, R.M., Van Dongen, B.F., Dumas, M., García-Bañuelos, L., Kunze, M., Leopold, H., Mendling, J., Uba, R., Weidlich, M., Weske, M., Yan, Z.: A Short Survey on Process Model Similarity. In: Bubenko, J., Krogstie, J., Pastor, O., Pernici, B., Rolland, C., Sølvberg, A. (eds.) *Seminal Contributions to Information Systems Engineering*, pp. 421–427. Springer Berlin Heidelberg, Berlin, Heidelberg (2013). https://doi.org/10.1007/978-3-642-36926-1_34, http://link.springer.com/10.1007/978-3-642-36926-1_34
7. Ganesh, C., Ramachandran, K., Varasree, B., Lakhanpal, S., Rohini, B., Acharjee, P.B.: Information Extraction Using Data Mining Techniques For Big Data Processing in Digital Marketing Platforms. In: *2023 10th IEEE Uttar Pradesh Section International Conference on Electrical, Electronics and Computer Engineering (UPCON)*. vol. 10, pp. 1662–1667. IEEE (2023)
8. Hoffmann, M., Bergmann, R.: Using graph embedding techniques in process-oriented case-based reasoning. *Algorithms* **15**(2) (2022). <https://doi.org/10.3390/a15020027>
9. Jurafsky, D., Martin, J.H.: *Speech and Language Processing*, <https://web.stanford.edu/~jurafsky/slp3/>
10. Kunze, M., Weidlich, M., Weske, M.: Behavioral Similarity – A Proper Metric. In: Rinderle-Ma, S., Toumani, F., Wolf, K. (eds.) *Business Process Management*, vol. 6896, pp. 166–181. Springer Berlin Heidelberg, Berlin, Heidelberg (2011). https://doi.org/10.1007/978-3-642-23059-2_15, http://link.springer.com/10.1007/978-3-642-23059-2_15, series Title: *Lecture Notes in Computer Science*
11. Le, Q., Mikolov, T.: Distributed representations of sentences and documents. In: *Proceedings of the 31st International Conference on Machine Learning. Proceedings of Machine Learning Research*, vol. 32, pp. 1188–1196. PMLR (2014)
12. Li, C., Reichert, M., Wombacher, A.: On measuring process model similarity based on high-level change operations. In: *International Conference on Conceptual Modeling*. pp. 248–264. Springer (2008)
13. McInnes, L., Healy, J., Astels, S., et al.: hdbscan: Hierarchical density based clustering. *J. Open Source Software* **2**(11), 205 (2017)
14. Mikolov, T., Chen, K., Corrado, G.S., Dean, J.: Efficient estimation of word representations in vector space (2013), <http://arxiv.org/abs/1301.3781>
15. Montani, S., Leonardi, G., Quaglini, S., Cavallini, A., Micieli, G.: A knowledge-intensive approach to process similarity calculation. *Expert Systems with Applications* **42**(9), 4207–4215 (Jun 2015). <https://doi.org/10.1016/j.eswa.2015.01.027>, <https://linkinghub.elsevier.com/retrieve/pii/S0957417415000421>, publisher: Elsevier

16. Murata, T.: Petri nets: Properties, analysis and applications. *Proceedings of the IEEE* **77**(4), 541–580 (1989). <https://doi.org/10.1109/5.24143>
17. Narayanan, A., Chandramohan, M., Venkatesan, R., Chen, L., Liu, Y., Jaiswal, S.: graph2vec: Learning distributed representations of graphs. *CoRR abs/1707.05005* (2017)
18. Peeva, V., van der Aalst, W.M.: Grouping Local Process Models. *arXiv preprint arXiv:2311.03040* (2023)
19. Pismerov, A., Pikalov, M.: Applying embedding methods to process mining. In: *Proceedings of the 5th International Conference on Algorithms, Computing and Artificial Intelligence. ACAI'22, Association for Computing Machinery* (2022). <https://doi.org/10.1145/3579654.3579730>
20. Recker, J., Rosemann, M., Indulska, M., Green, P.: Business process modeling—a comparative analysis. *Journal of the association for information systems* **10**(4), 1 (2009)
21. Řehůřek, R., Sojka, P.: Software framework for topic modelling with large corpora. In: *Proceedings of the LREC 2010 Workshop on New Challenges for NLP Frameworks*. pp. 45–50. ELRA (May 2010)
22. Syukriilah, N., Kusumo, D.S., Widowati, S.: Structural similarity analysis of business process model using selective reduce based on Petri Net. In: *2015 3rd International Conference on Information and Communication Technology (ICoICT)*. pp. 1–5. IEEE (2015)
23. Sánchez-Charles, D., Muntés-Mulero, V., Carmona, J., Solé, M.: Process model comparison based on cophenetic distance. In: *Business Process Management Forum: BPM Forum 2016, Rio de Janeiro, Brazil, September 18–22, 2016, Proceedings 14*. pp. 141–158. Springer (2016)
24. Van Der Aalst, W.M.P., De Medeiros, A.K.A., Weijters, A.J.M.M.: Process Equivalence: Comparing Two Process Models Based on Observed Behavior. In: Hutchison, D., Kanade, T., Kittler, J., Kleinberg, J.M., Mattern, F., Mitchell, J.C., Naor, M., Nierstrasz, O., Pandu Rangan, C., Steffen, B., Sudan, M., Terzopoulos, D., Tygar, D., Vardi, M.Y., Weikum, G., Dustdar, S., Fiadeiro, J.L., Sheth, A.P. (eds.) *Business Process Management*, vol. 4102, pp. 129–144. Springer Berlin Heidelberg, Berlin, Heidelberg (2006). https://doi.org/10.1007/11841760_10, http://link.springer.com/10.1007/11841760_10, series Title: Lecture Notes in Computer Science
25. Van Dongen, B., Dijkman, R., Mendling, J.: Measuring similarity between business process models. *Seminal Contributions to Information Systems Engineering: 25 Years of CAiSE* pp. 405–419 (2013), publisher: Springer
26. Verbeek, E.: Process discovery contest (2023), <https://icpmconference.org/2023/process-discovery-contest/>
27. Yan, Z., Dijkman, R., Grefen, P.: Fast business process similarity search. *Distributed and Parallel Databases* **30**, 105–144 (2012), publisher: Springer
28. Zeng, Z., Tung, A., Wang, J., Feng, J., Zhou, L.: Comparing Stars: On Approximating Graph Edit Distance. *PVLDB* **2**, 25–36 (Jan 2009)
29. Zhou, C., Liu, C., Zeng, Q., Lin, Z., Duan, H.: A Comprehensive Process Similarity Measure Based on Models and Logs. *IEEE Access* **7**, 69257–69273 (2019). <https://doi.org/10.1109/ACCESS.2018.2885819>, <https://ieeexplore.ieee.org/document/8611068/>