

Temporal Pattern Discovery for Scalable and Interpretable Time Series Analysis



Len Feremans
Hasselt University
Data Science Institute



Some research questions

- Goal is to find and report interesting patterns towards end-users
 - What patterns are interesting?
 - How can we find interesting patterns efficiently?
 - How can we use patterns to make more accurate predictions?

Outline

1. Pattern Mining

2. Representation

- Interpretable analysis of time series using patterns embeddings

3. Motif discovery

- Finding motifs in time series

Pattern Mining

- What, why, how of pattern mining?
- Applications

Why pattern mining?

- **Exploratory data analysis:** Gems in every dataset waiting to be discovered
- **Extract interpretable features**
 - Intrinsic interpretable representations and models
 - Using rules as post-hoc explainable model
- **Human in the loop discovery**
 - Add constraints and domain knowledge
- **Large amount of prior research**
 - Hundreds of efficient algorithms
 - See <https://www.philippe-fournier-viger.com/spmf/>



Algorithms

SPMF offers implementations of the following data mining algorithms.

🔍 Search algorithms, e.g. FP-Growth, HUI, clustering...

Sequential Pattern Mining

These algorithms discover sequential patterns in a set of sequences. For a good overview of **sequential pattern mining**, see this [paper](#).

What is Pattern?

- An **itemset**: $X=\{A,B\}$
 - Examples where both A and B co-occur
- A **sequential pattern**: $X=[A,B]$
 - Examples where A occurs before B
- A **rule**: $X = A \rightarrow B$
 - If A occurs in an example, B also occurs with a certain likelihood
- A time series **motif**: $X=(A,B)$
 - A set of subsequences with high similarity

Properties of patterns

- **Large search space**

- Singletons: {A},{B},... N
- Pairs: {A,B}, {B,C}, {A,C} $\binom{N}{2}$
- Triples: {A,B,C} $\binom{N}{3}$
- Possible itemsets 2^N

- **Monotonicity**

- If {A,B,C} occurs in **S** examples then any subset occurs at least **S** times

- **Constraints**

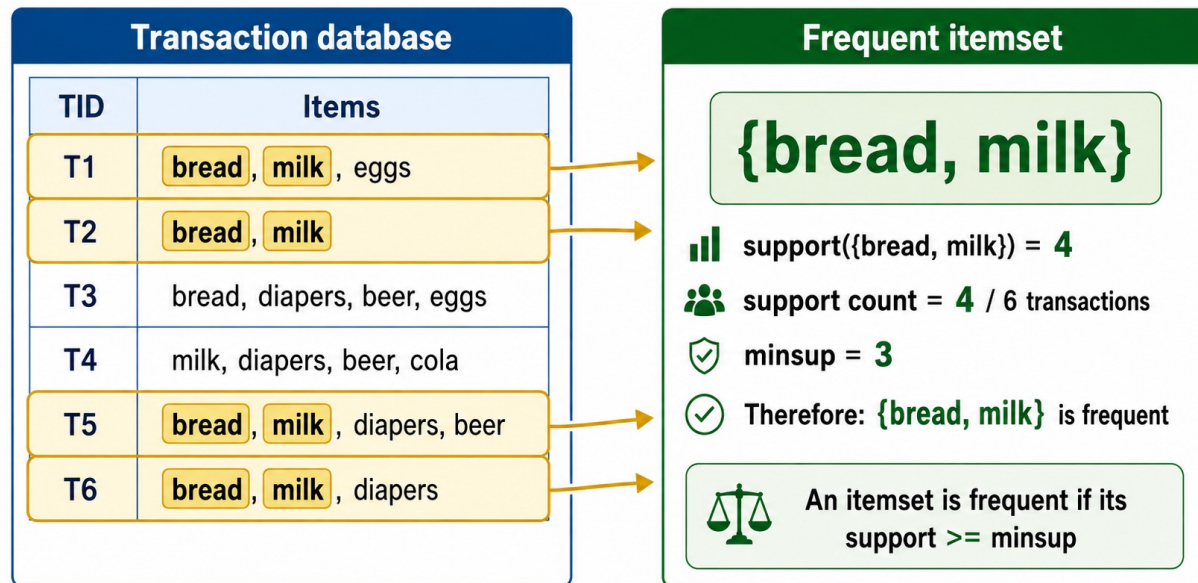
- Limit on size, e.g. k=3
- Limit on number of patterns, e.g. top-k=10.000
- Gap constraints

- **Value of pattern**

- Frequency
- Confidence
- Pos. or neg. for label
- Cohesive

Applications

- **Frequent Itemset Mining** for Market-basket analysis [Agrawal, 1994]



Applications

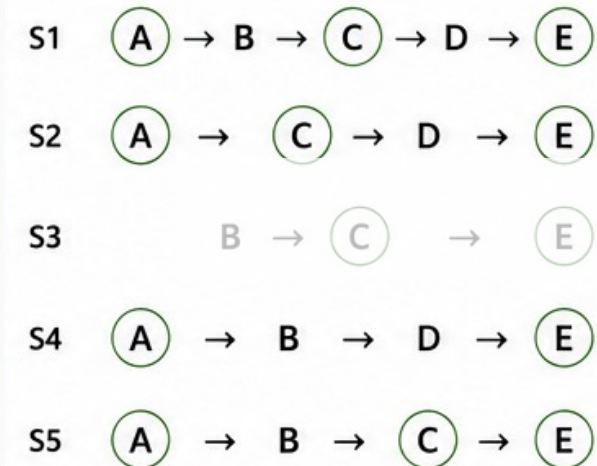
- Sequential pattern mining

1. Input: Sequence Database

Session ID	Sequence of items (in time order)
S1	A → B → C → D → E
S2	A → C → D → E
S3	B → C → E
S4	A → B → D → E
S5	A → B → C → E

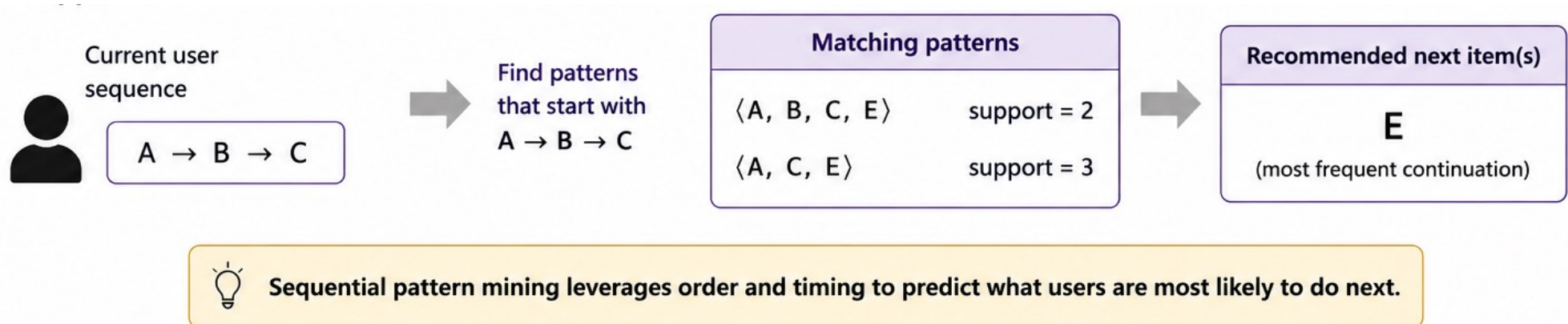
Items: A, B, C, D, E (e.g., songs, videos, products)

Example pattern: $\langle A, C, E \rangle$



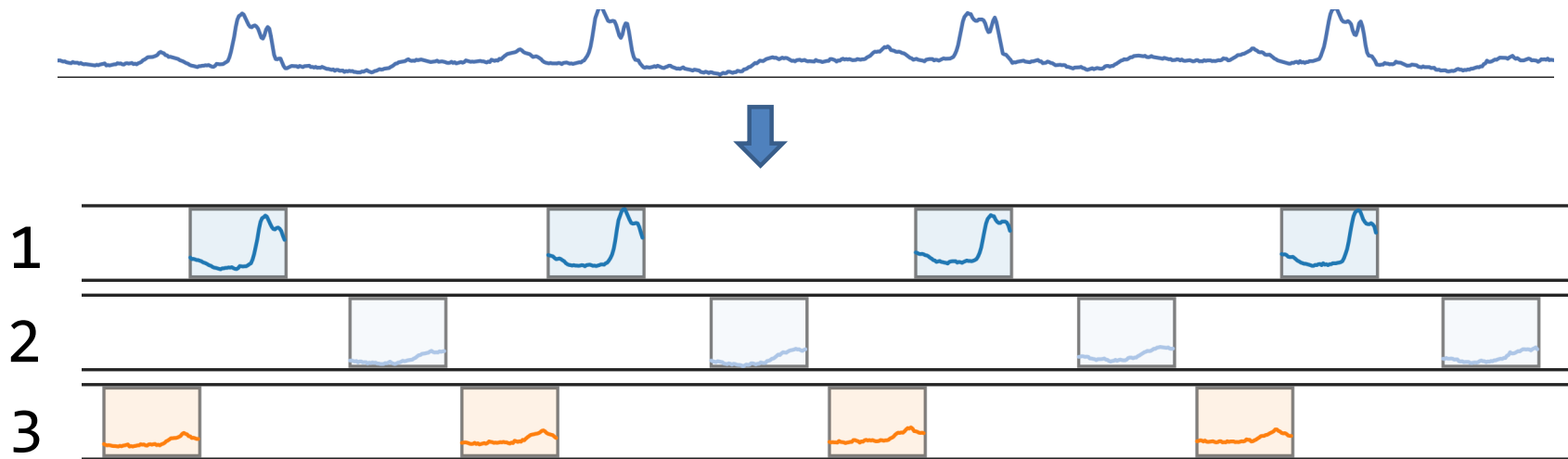
Applications

- Recommender systems



Applications

- Time series motifs
 - Repeating (non-overlapping) subsequences with high similarity



Applications



- Time series (heart rate, blood pressure, oxygen)
 - Events (alarms)



Rule mining



Insights



Smart alarms

(e)	$\{HF < 60, 28 < RF < 31\} \rightarrow \{intubation\}$	0.55	0.02
(f)	$\{74 < HF < 86, RF > 30\} \rightarrow \{intubation\}$	0.53	0.03
(g)	$\{ABP_d < 47, RF > 30\} \rightarrow \{intubation\}$	0.52	0.05
(h)	$\{ABP_d < 45, ABP_m < 64, RF > 31\} \rightarrow \{intubation\}$	0.51	0.02
(i)	$\{intubation, ABP_d < 45, 91 < S_pO_2 < 93\} \rightarrow \{ABP_m < 64\}$	0.51	0.02

HF < 60 and
28 < RF < 32? 55% chance of intubation

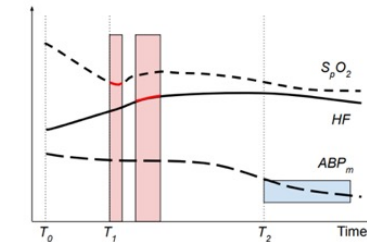


Fig. 1: Example of multivariate time series of a patient with alarms. S_pO_2 oxygen saturation, HF heart frequency, ABP_m arterial blood pressure.

Occurs in 0.2% of patients
(+ 15 patients) [Venturini, 2025]

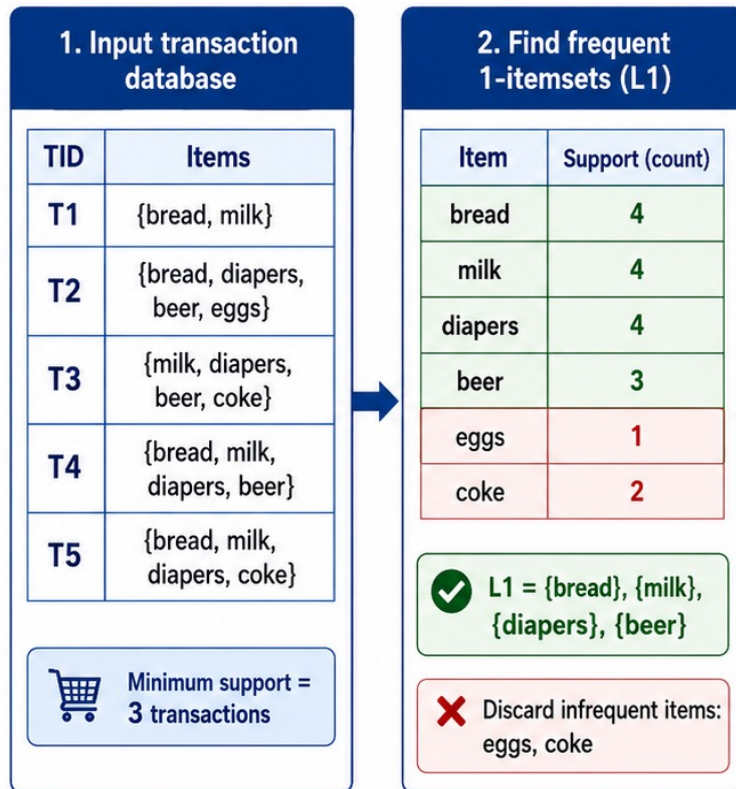
Apriori algorithm for frequent itemset mining

1. Input transaction database	
TID	Items
T1	{bread, milk}
T2	{bread, diapers, beer, eggs}
T3	{milk, diapers, beer, coke}
T4	{bread, milk, diapers, beer}
T5	{bread, milk, diapers, coke}

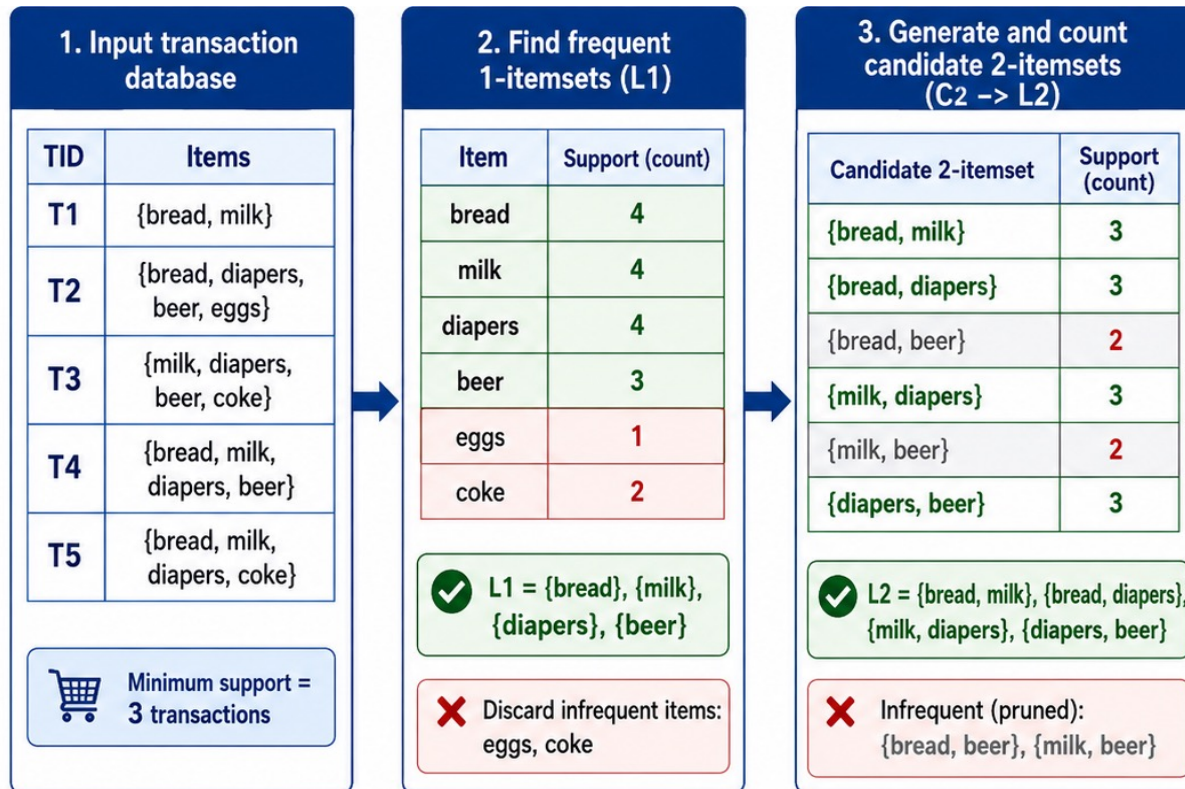
➔

 Minimum support = 3 transactions

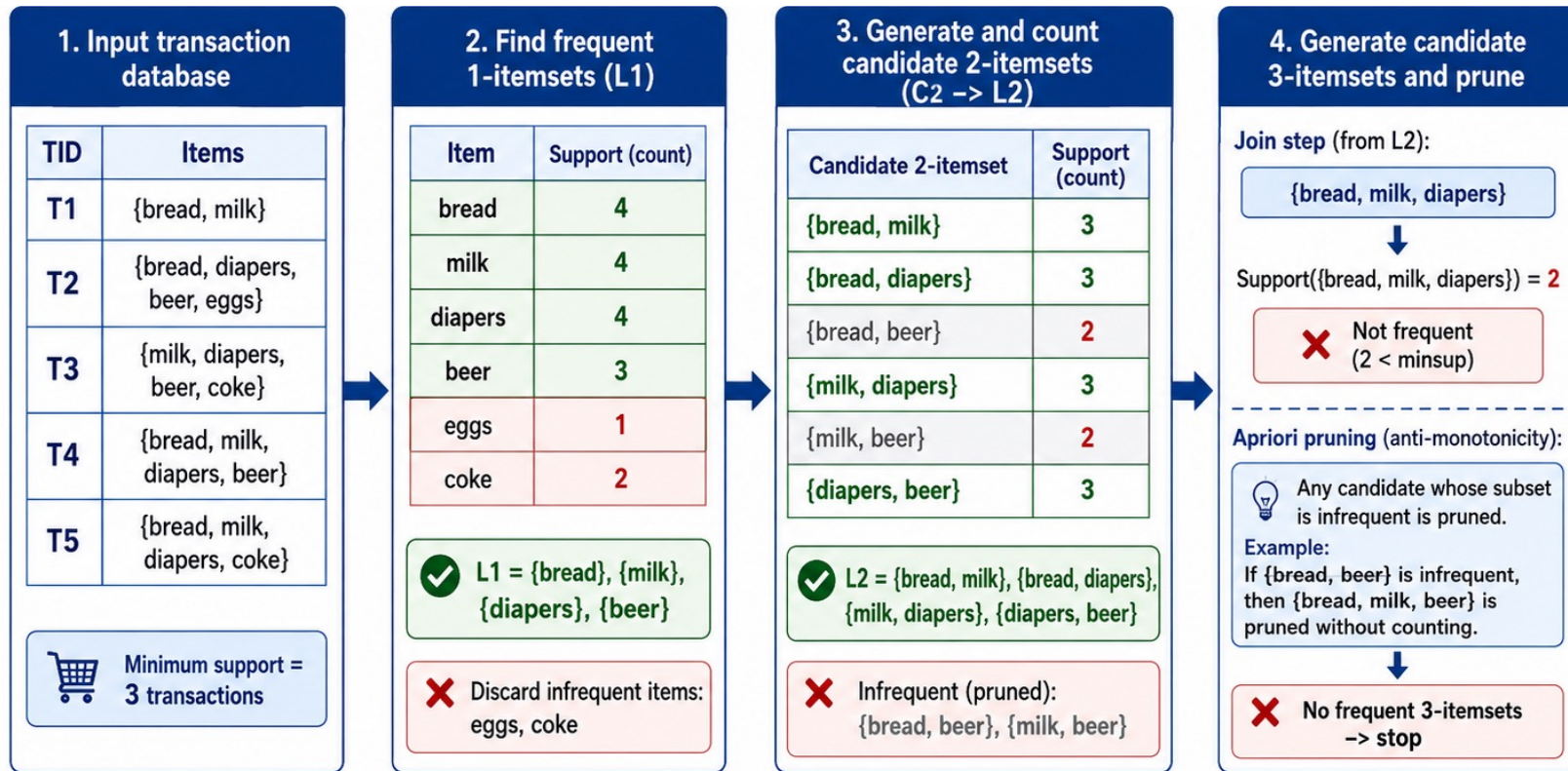
Apriori algorithm for frequent itemset mining



Apriori algorithm for frequent itemset mining



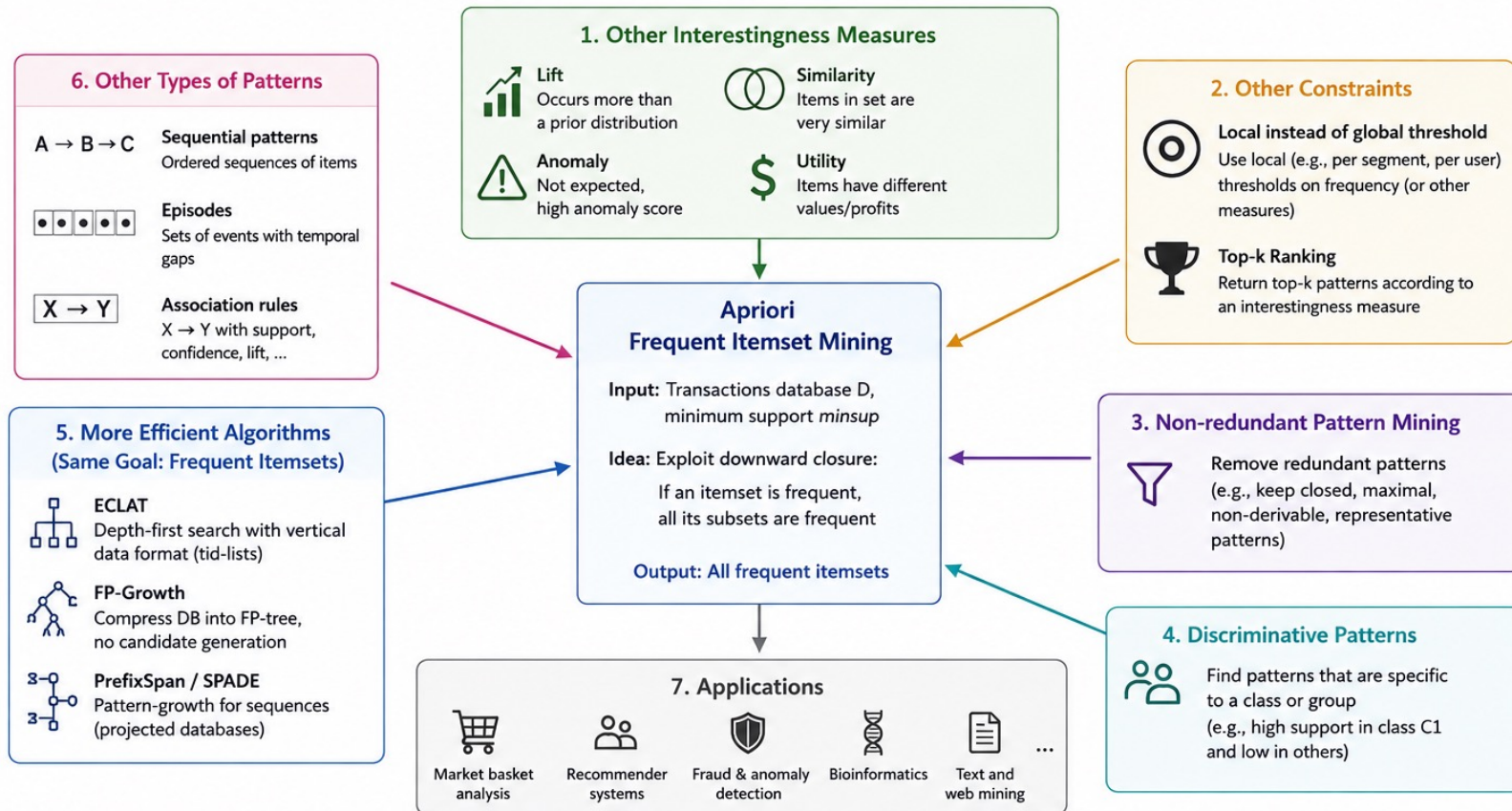
Apriori algorithm for frequent itemset mining



Apriori counts candidates level by level and uses:

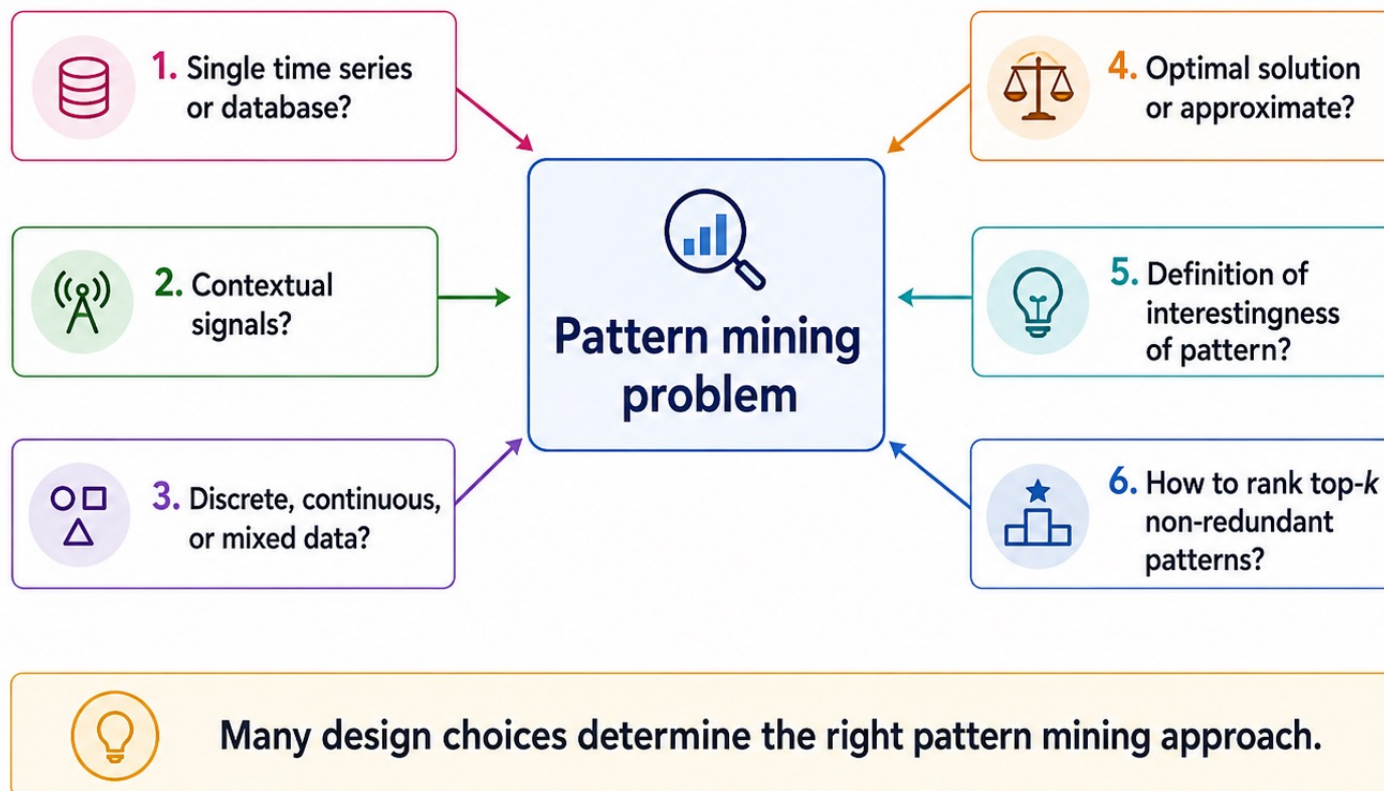
if an itemset is infrequent, all of its supersets are infrequent.

Extensions to Apriori



Apriori is the foundation. Research extends it by changing interestingness, constraints, pattern types, redundancy handling, discriminative power, and efficiency.

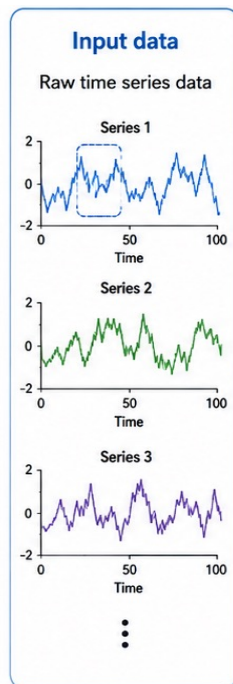
Bigger picture: pattern mining problems



Interpretable and scalable analysis of TS using **patterns embeddings**

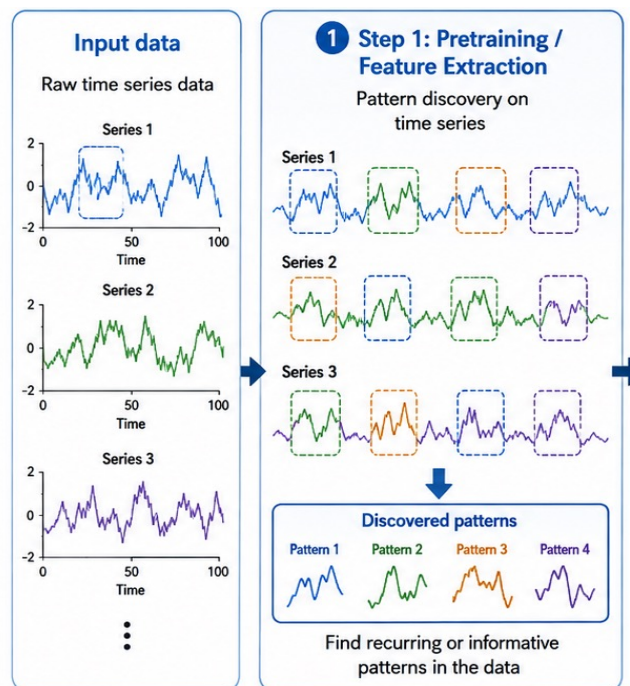
- **PBAD** [Feremans, 2019]:
 - **Anomaly detection** in mixed-type event sequences
- **PETSC** [Feremans, 2022]
 - **Time Series Classification**

What are pattern-based embeddings



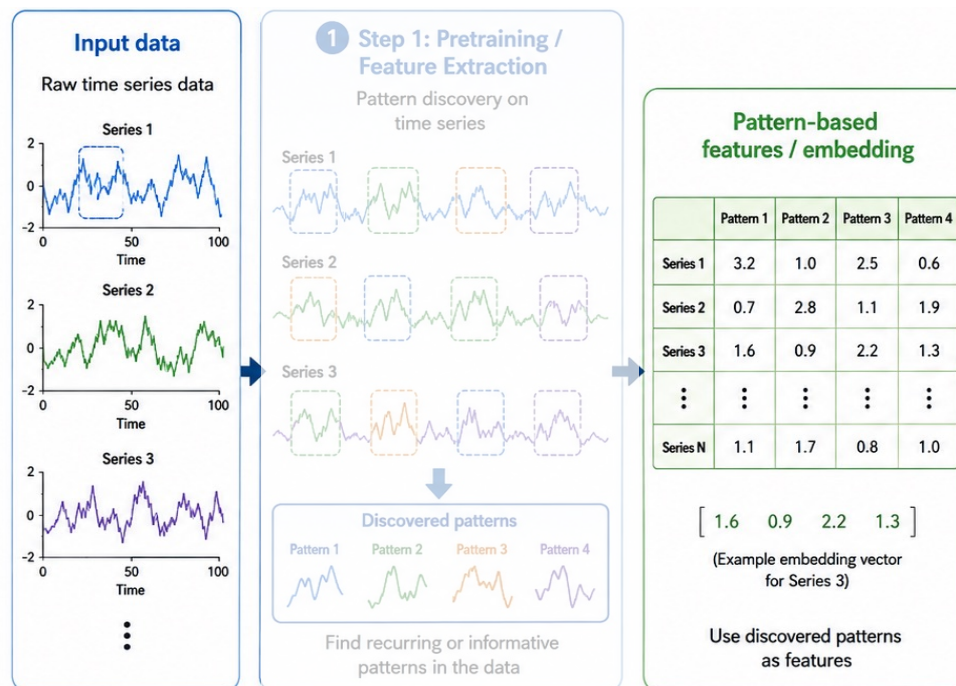
First discover patterns, then use them as features for prediction.

What are pattern-based embeddings



First discover patterns, then use them as features for prediction.

What are pattern-based embeddings



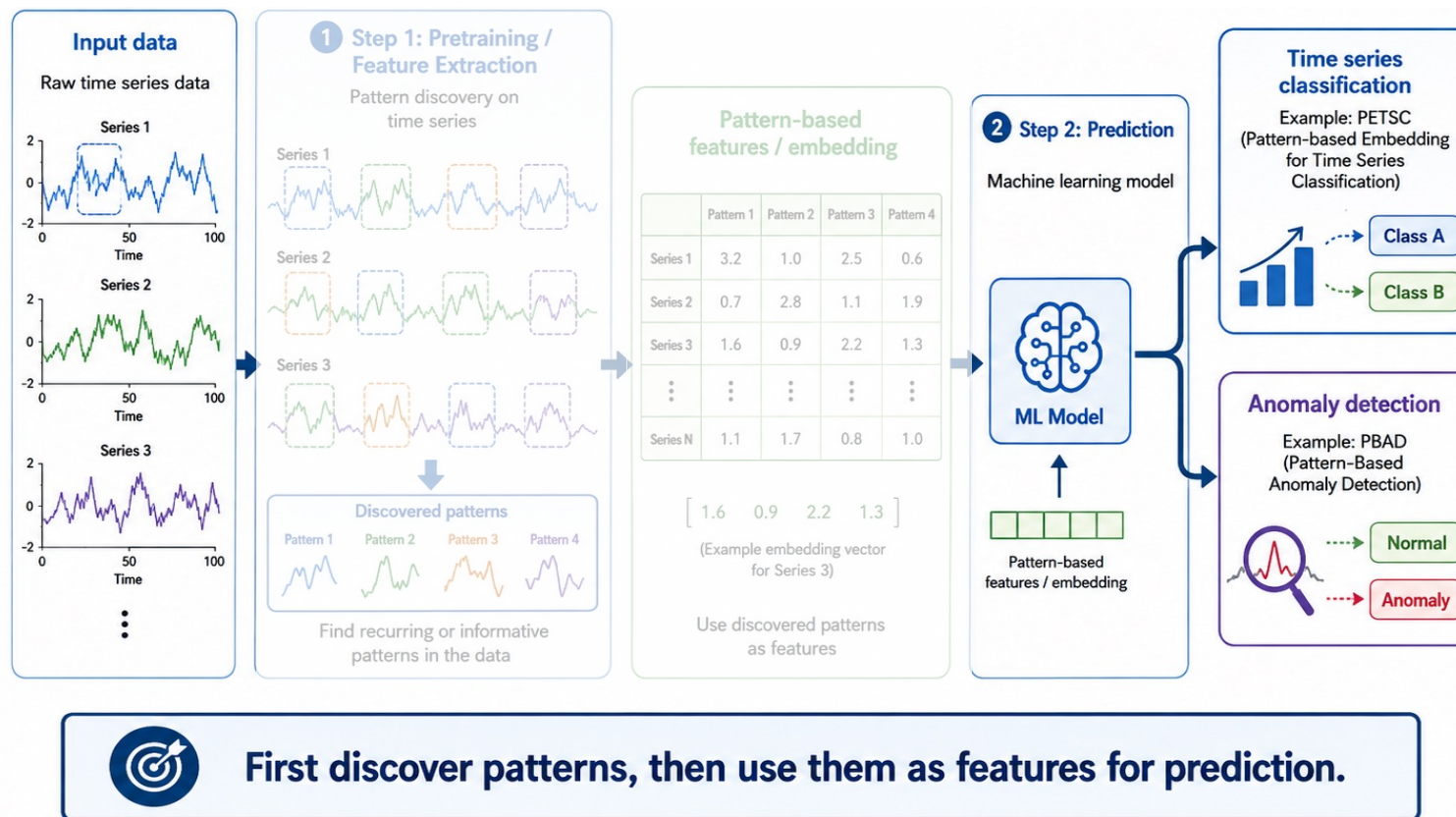
First discover patterns, then use them as features for prediction.

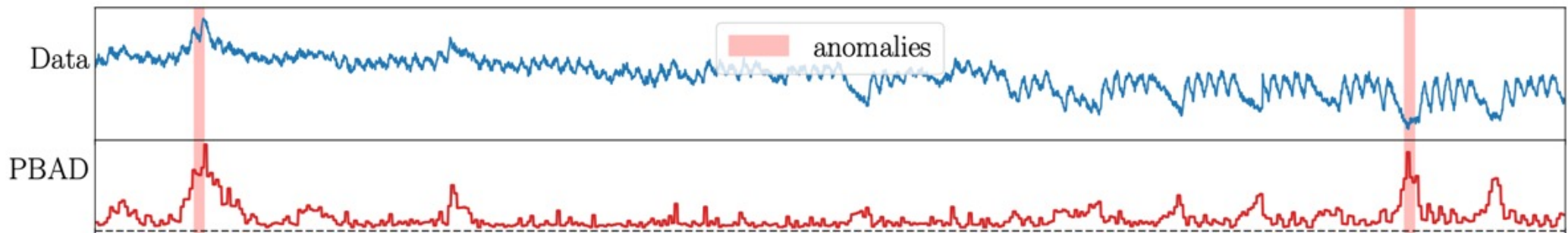
What are pattern-based embeddings



First discover patterns, then use them as features for prediction.

What are pattern-based embeddings

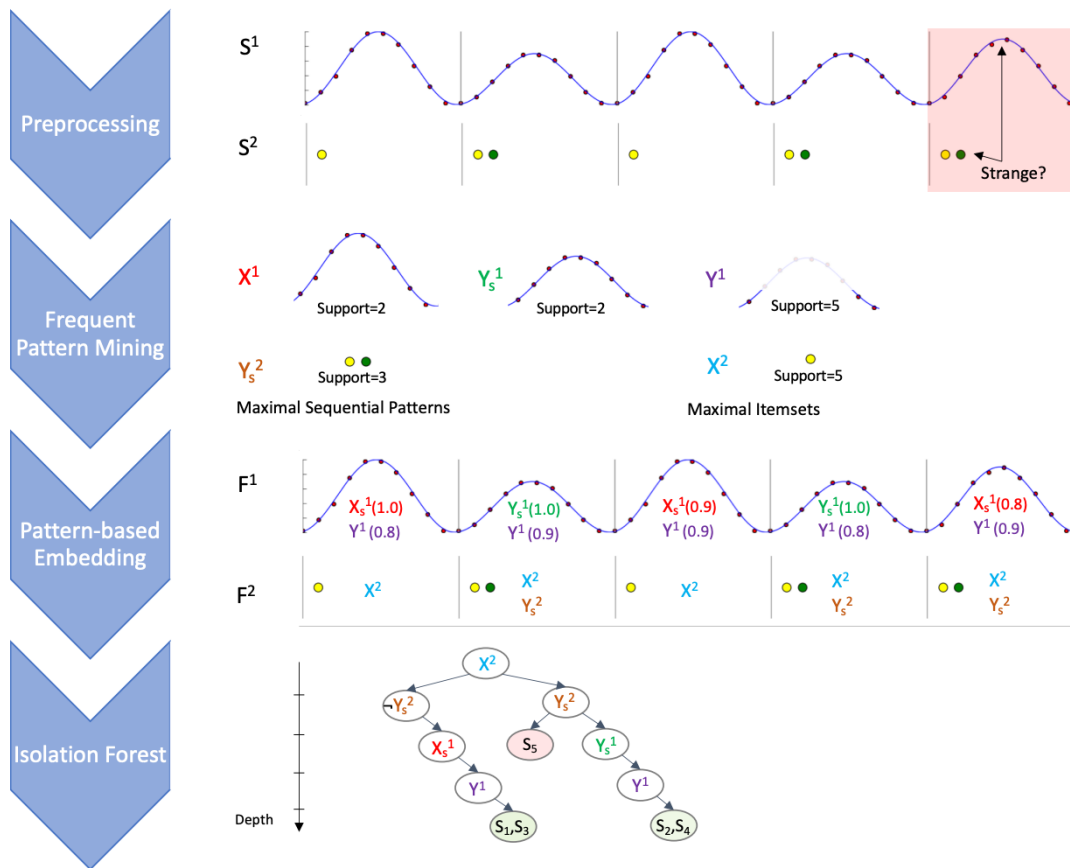




- If at window X_t , **few frequent patterns occur**, there is a high likelihood that there is an **anomaly** at X_t

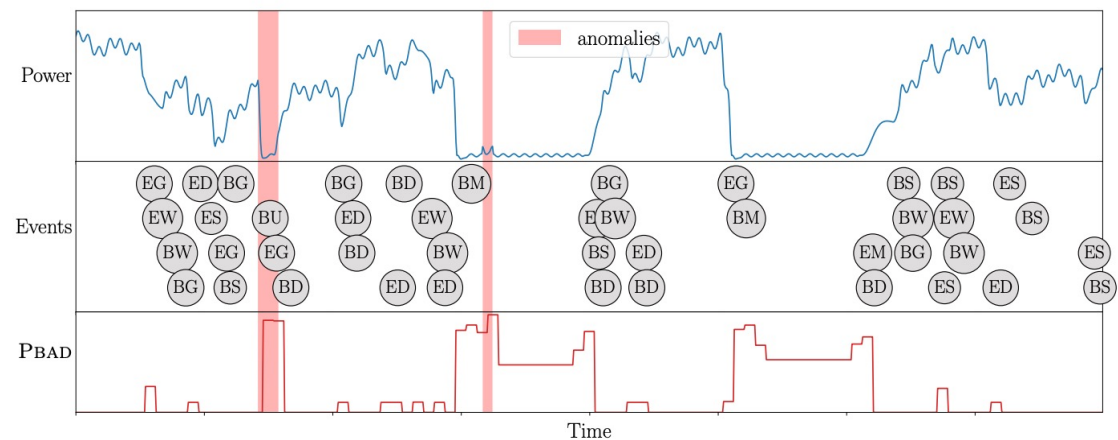


Pattern-based anomaly detection in mixed-type time series



Pattern-based anomaly detection in mixed-type time series

- Unlike compared baselines, PBAD jointly models continuous time series and event data. By combining patterns with **contextual features**, such as events, alarms or multivariate signals



- In our 2019 study, PBAD outperform time series anomaly detection techniques, including MatrixProfile, Pav, Mifpod and FPOF.

TIPM: Pattern mining and anomaly detection in multi-dimensional time series and event logs

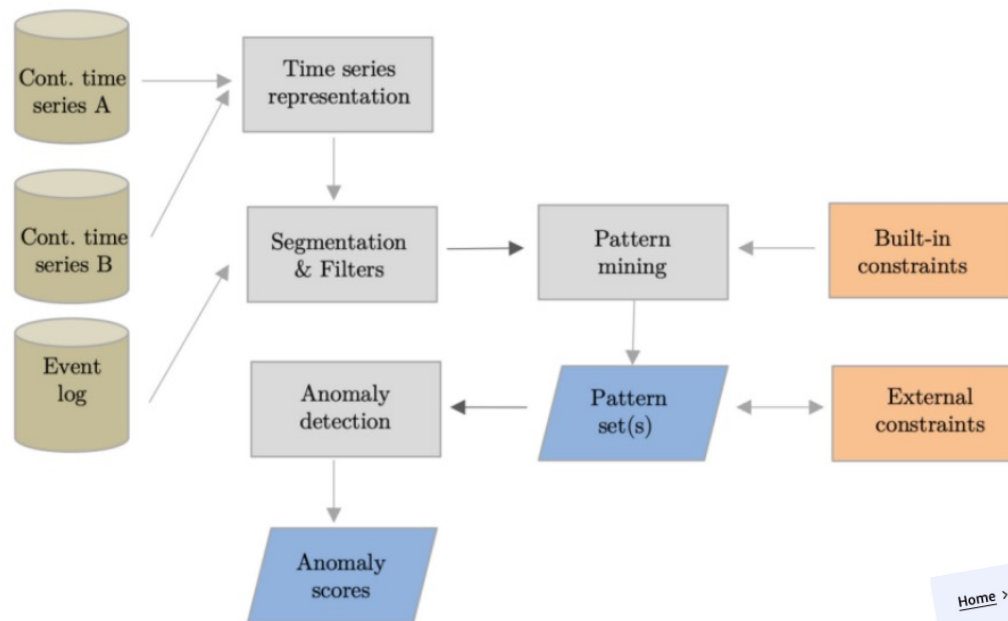
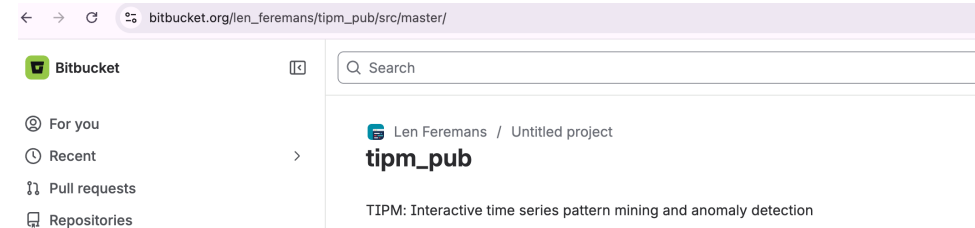
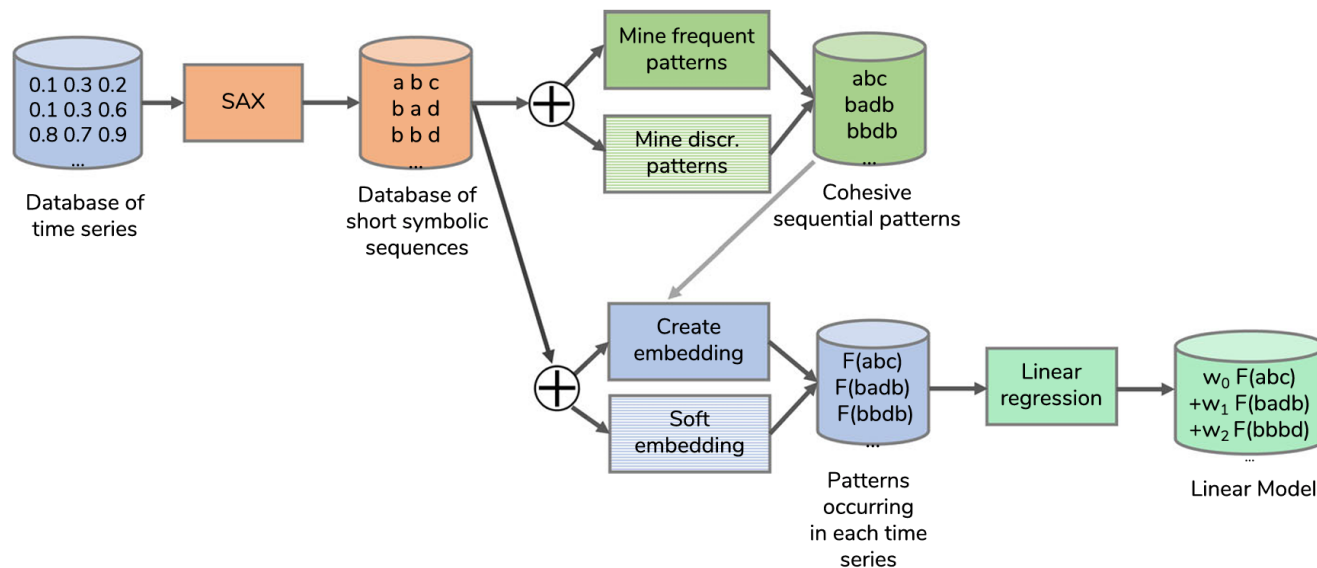


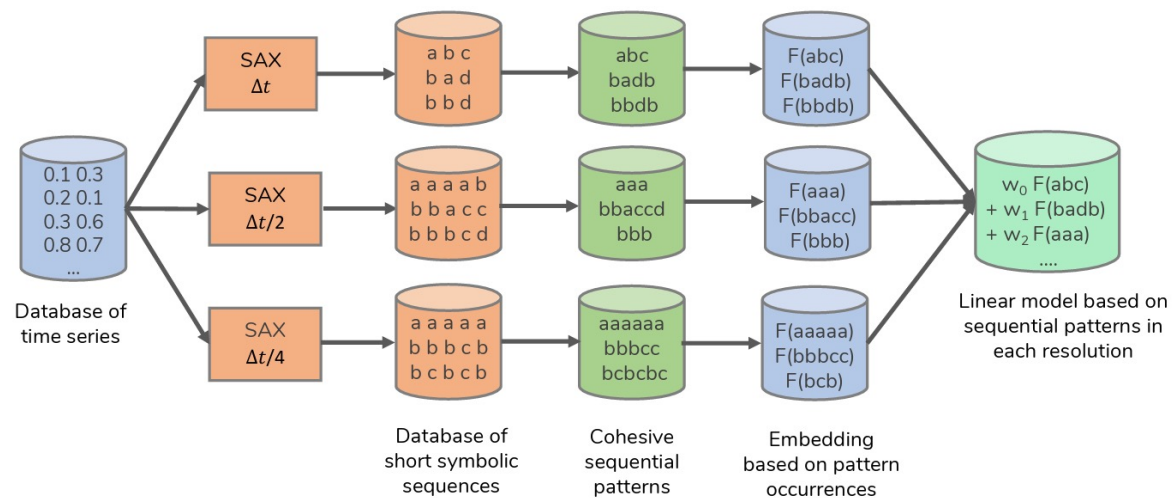
Fig. 2. Workflow of our framework.



PETSC: Pattern-based embedding for time series classification



PETSC: Pattern-based embedding for time series classification

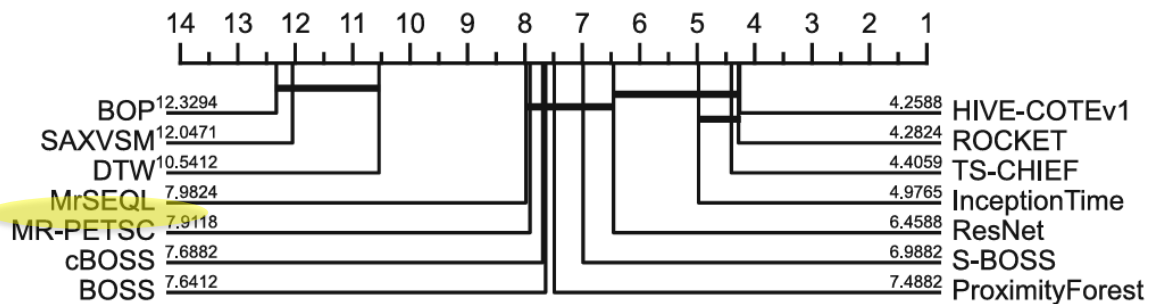
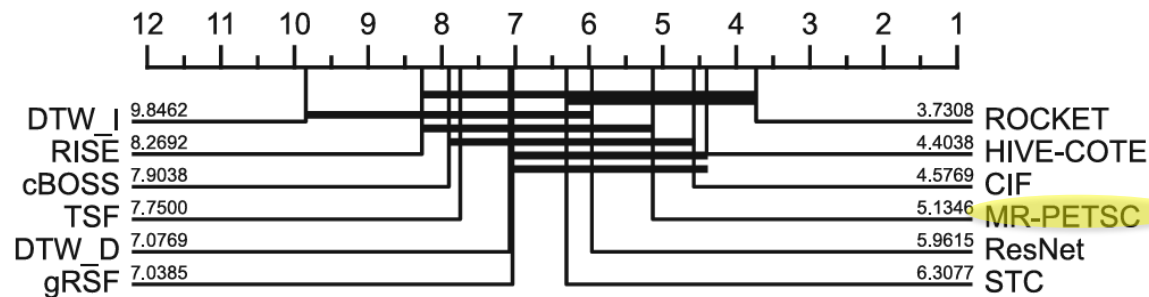


An ensemble combining **different resolutions** works best

PETSC: Pattern-based embedding for time series classification

- Competitive accuracy compared to deep learning and ensemble learning (2022)

26 multivariate datasets



85 univariate datasets

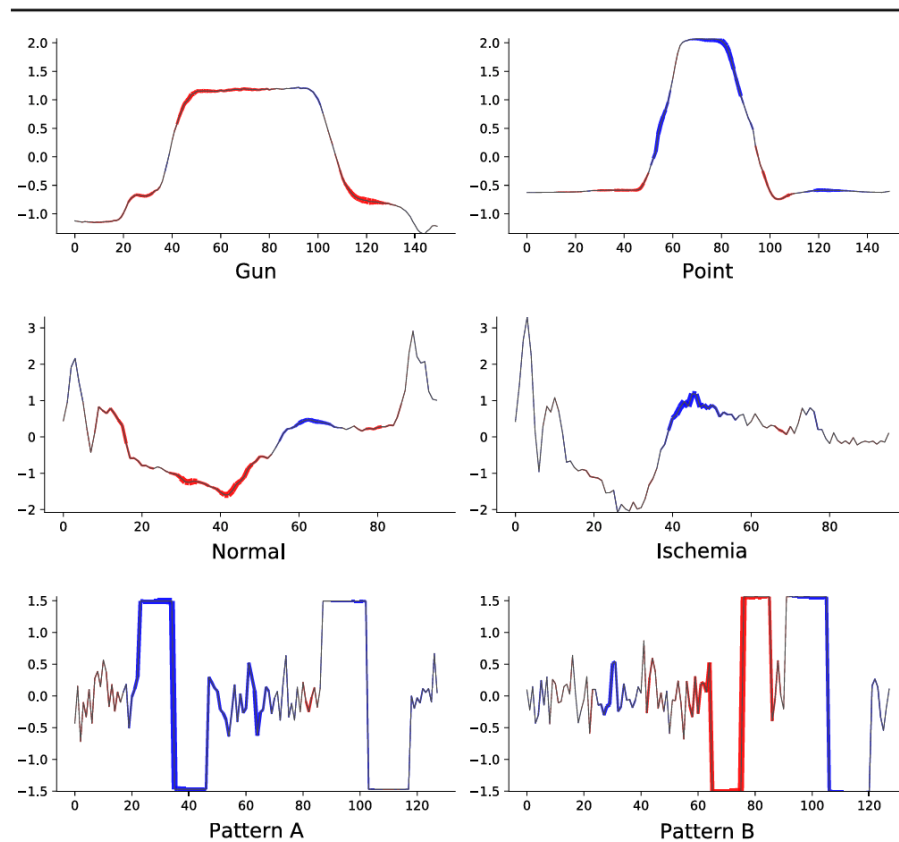
PETSC: Pattern-based embedding for time series classification

- Order-of-magnitude faster compared to most deep learning and ensemble methods

	ROCKET	MR- PETSC	cBOSS	MiSTiCL	MR- SEQL	INCEPTION	TIME PROXIMITY	FOREST TS- CHIEF
<i>Total time</i>	1.4h	<u>2.7h</u>	19.5h	20.2h	23.7h	6d	11d	11d

Total time for **training and predictions**
on **85** univariate datasets

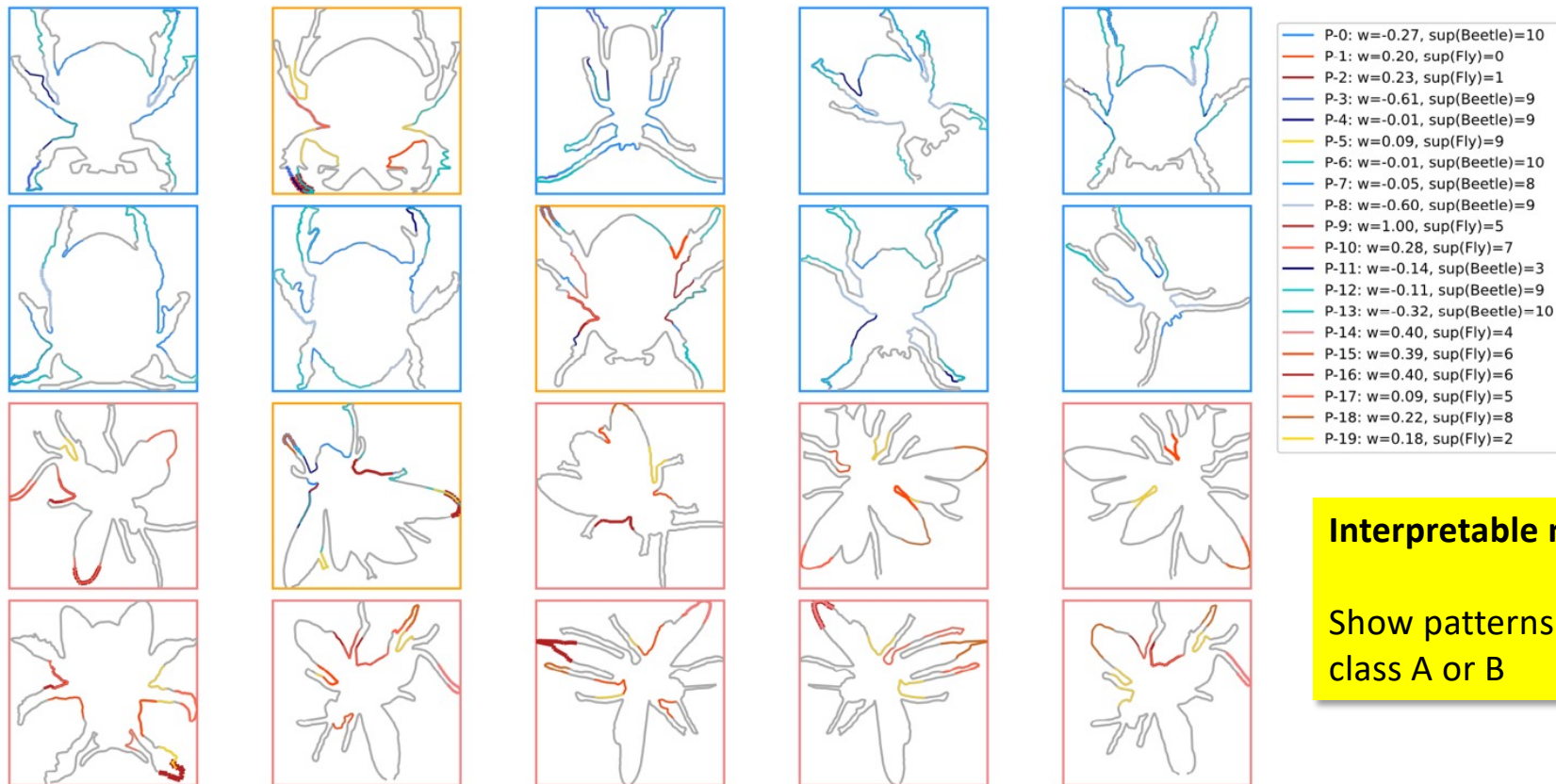
PETSC: Pattern-based embedding for time series classification



**Interpretable decisions
case-by-case!**

Highlight part of time series
typical of class A or B

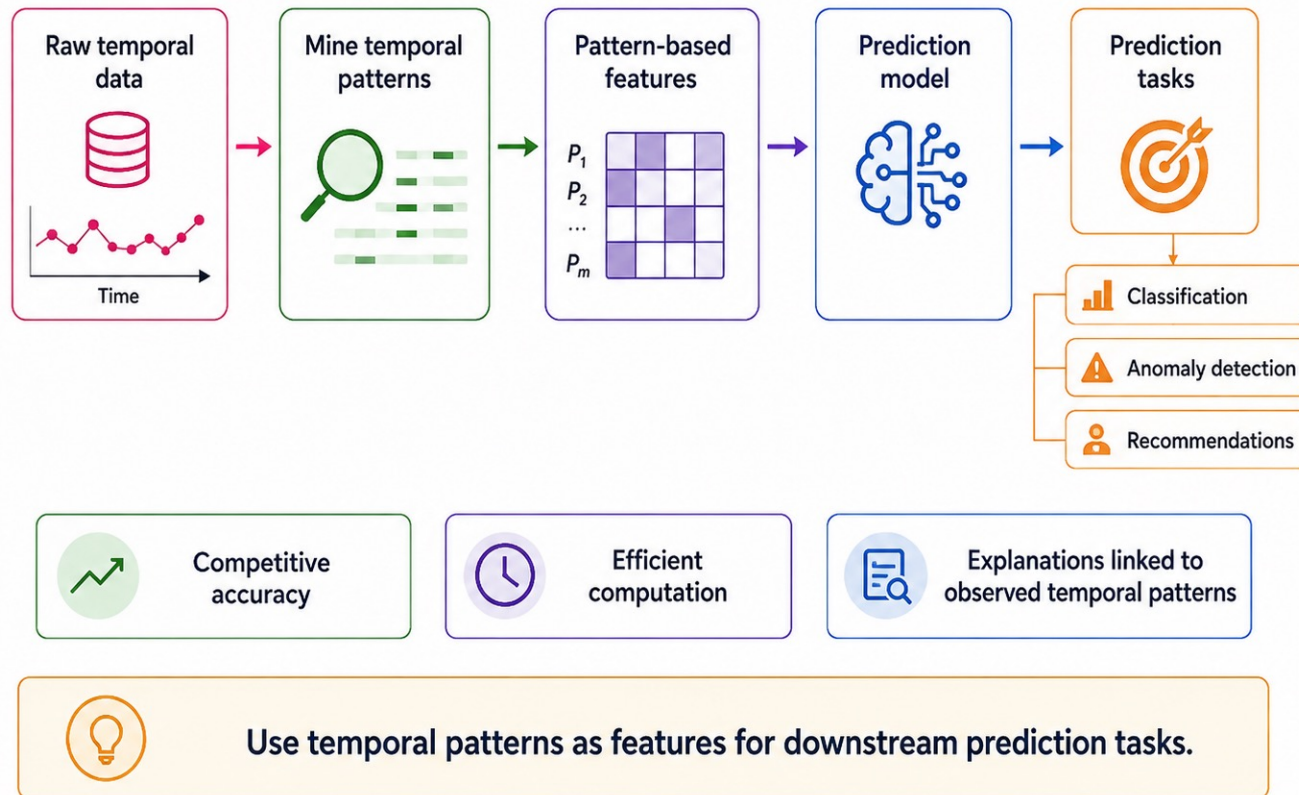
PETSC: Pattern-based embedding for time series classification



Interpretable model!

Show patterns typical for class A or B

Temporal patterns for prediction



Finding motifs in time series

- Find the **best** motifs approximate and exact
- Find the **best collection** of motifs using MotiPlus and MotiSet

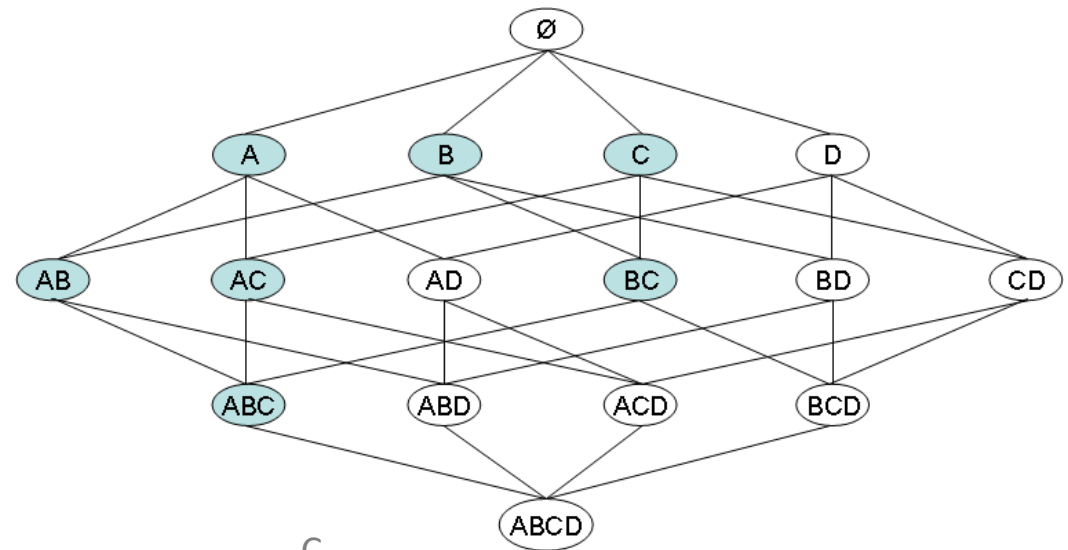
Motivation

- Existing motif discovery methods

- ✗ Use **heuristics**
- ✗ Lower **accuracy** on top-k discovery
- ✗ **Sensitive parameters**
- ✓ Quadratic **runtime** performance.

Motivation

- Motif mining
 - Variant of Apriori
- Motif set mining
 - Variant of Clustering



Motivation

- We want to **discover** the **top-k** most interesting motifs
 - Example: Lyrics and instrumental sections in music

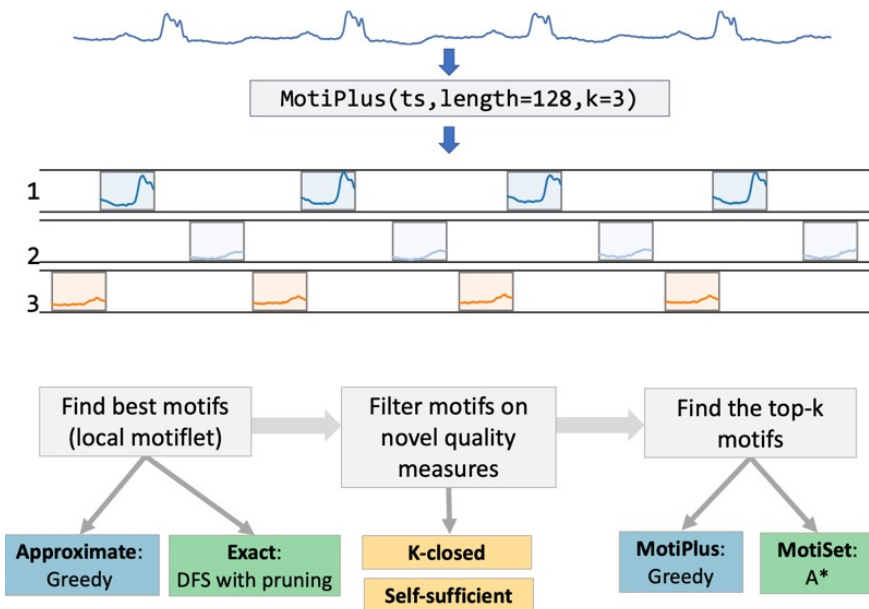


MotiPlus and MotiSet: Discovering the Best Set of Motiflets in Time Series

Len Feremans, Adrem Data Lab, University of Antwerp, Belgium
Patrick Schäfer Humboldt-Universität zu Berlin, Germany
Wannes Meert DTAI Research Group, KU Leuven, Belgium

bitbucket.org/len_feremans/kmotiflets

Motif discovery: Find repeating subsequences with high similarity



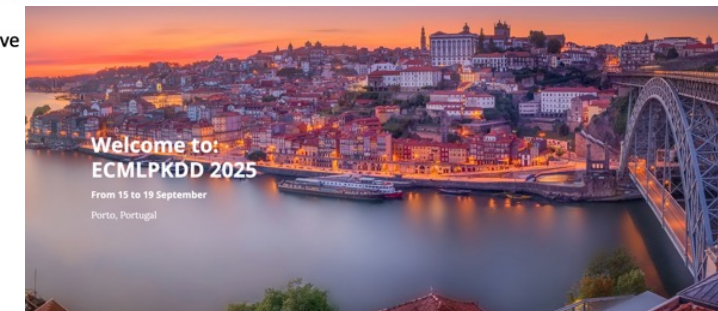
Existing approaches **have limitations**

1. Existing motif discovery methods such as GrammarViz, MMotif, Motiflets and LoCoMotif:

- ✗ Have lower **precision** and **recall** on the **TSMD-benchmark**.
- ✗ Performance varies widely with **sensitive parameters** such as the distance threshold.
- ✗ Uses **heuristics** for discovering motifs and collections of motifs.
- ✓ Have excellent **runtime** performance.

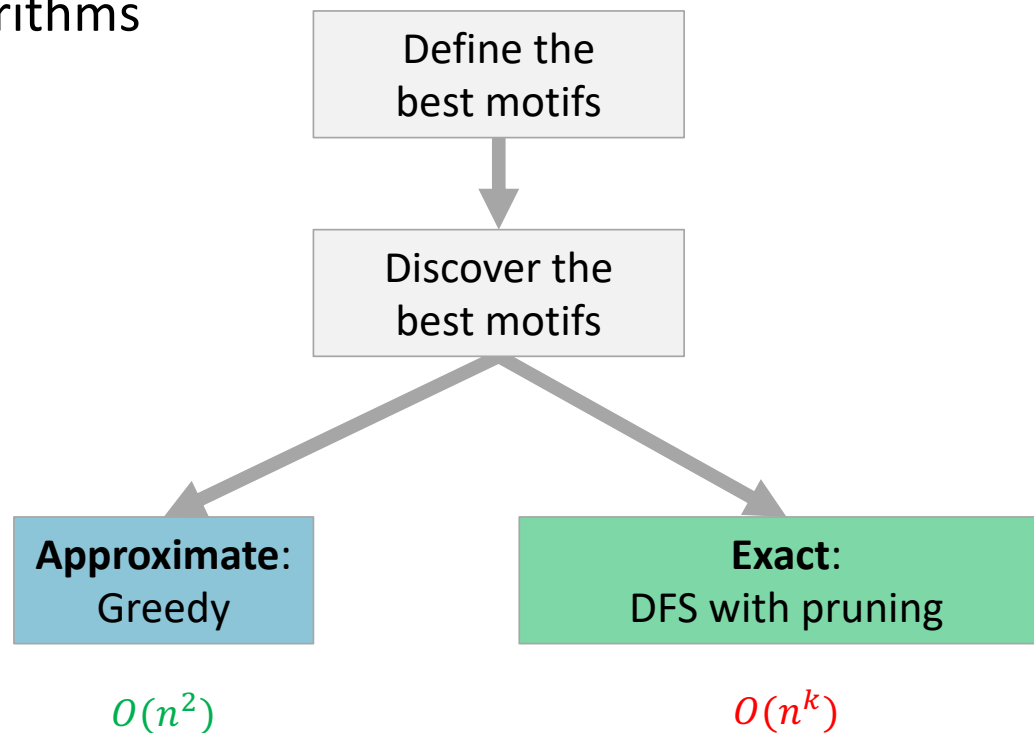
2. **MotiPlus** and **MotiSet** are novel algorithms that discover a set of motifs

- ✓ Have increased **precision** (+27%) and **recall** (+10%) on the **TSMD-benchmark**.
- ✓ Easy-to-tune **stable parameters**, i.e. maximal size of motif and motif collection.
- ✓ **A* search** and **exact motif search** guarantee **optimal** motifs and sets thereof.
- ✓ Greedy and anytime A* algorithms have excellent **runtime** performance.

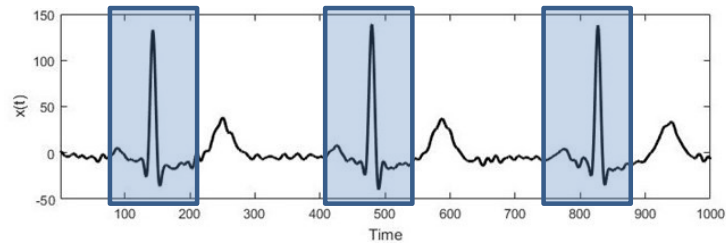
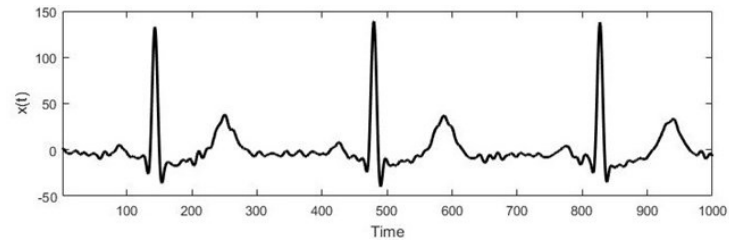


First contribution

- Novel definition of **the best motif** and **quality constraints**, i.e. **local motiflets**
- **Discovery** algorithms



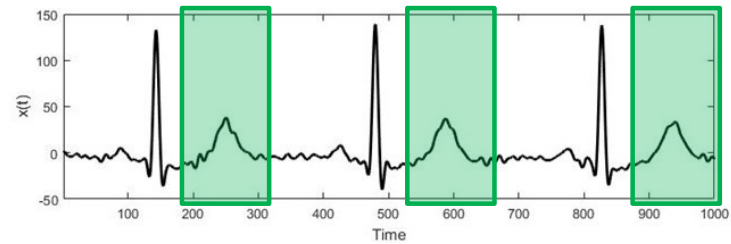
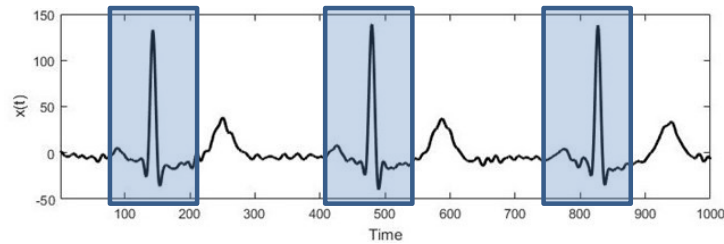
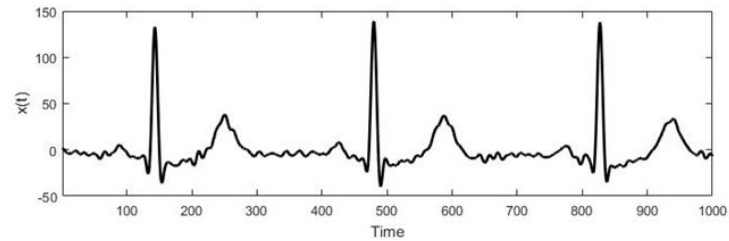
Find **k-motiflets**
[Schäfer and Leser,
2022]



k-motiflet: Candidate motif with k non-overlapping subsequences with highest similarity*

*measured using maximal pairwise z-normalized euclidean distance

Find **local motiflets**



Local k-motiflet: Rank on similarity and not **redundant** to other local k-motiflets

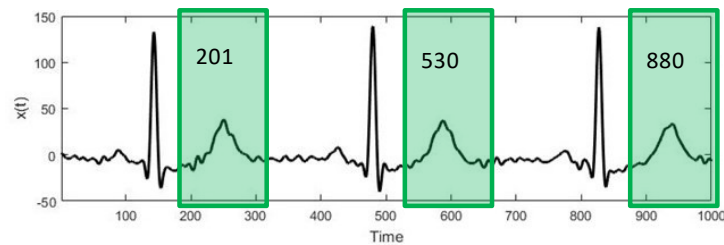
Filter motifs on novel
quality measures



K-closed

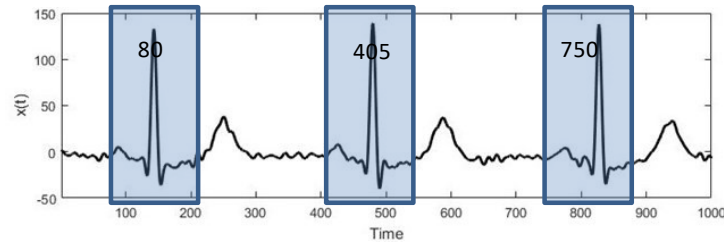
Self-sufficient [Webb, 2010]

Search for motiflets



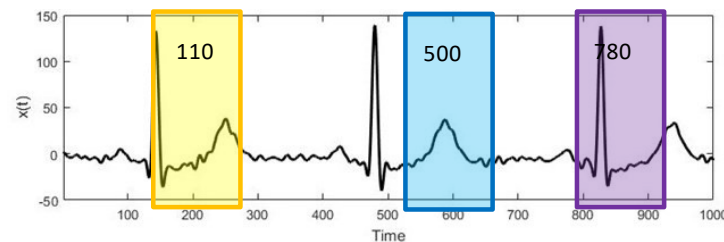
Distance:
0.5

Approximate search



Distance:
0.3

Exact search

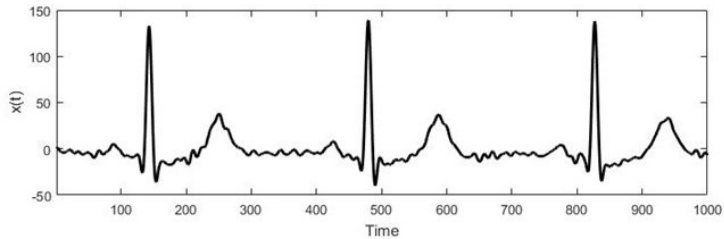


Distance: 5

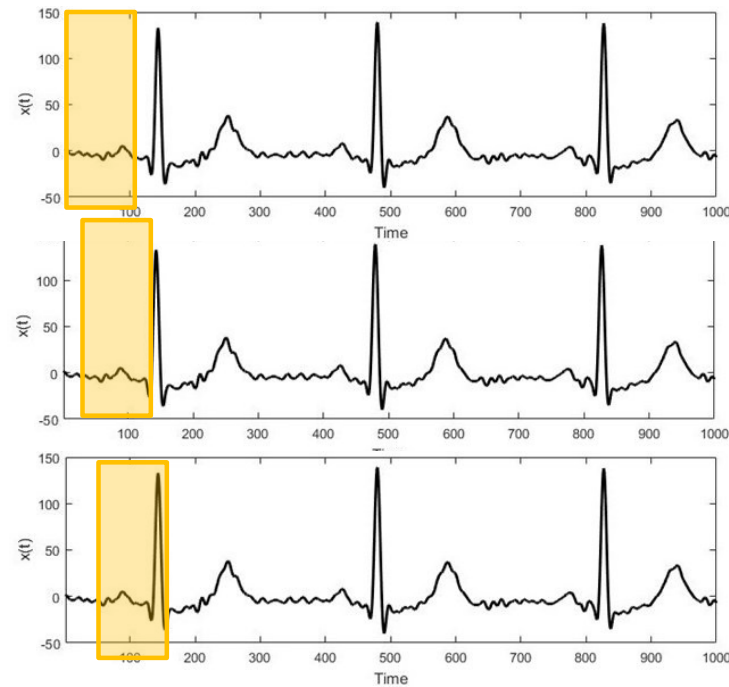
Monotonicity: If motif $\{A,B\}$ has a high distance, any superset of $\{A,B\}$ will also have a high distance

Search for motiflets

Motif discovery is **discrete**, there are $N - \text{window_length} + 1$ candidate subsequence



Sliding window



1

2

3

...

Next, we compute the **similarities** between all subsequence pairs

Time Series

X1	X2	X3	X4	X5	X6	X7	X8	X9	X10	X11	X12
----	----	----	----	----	----	----	----	----	-----	-----	-----

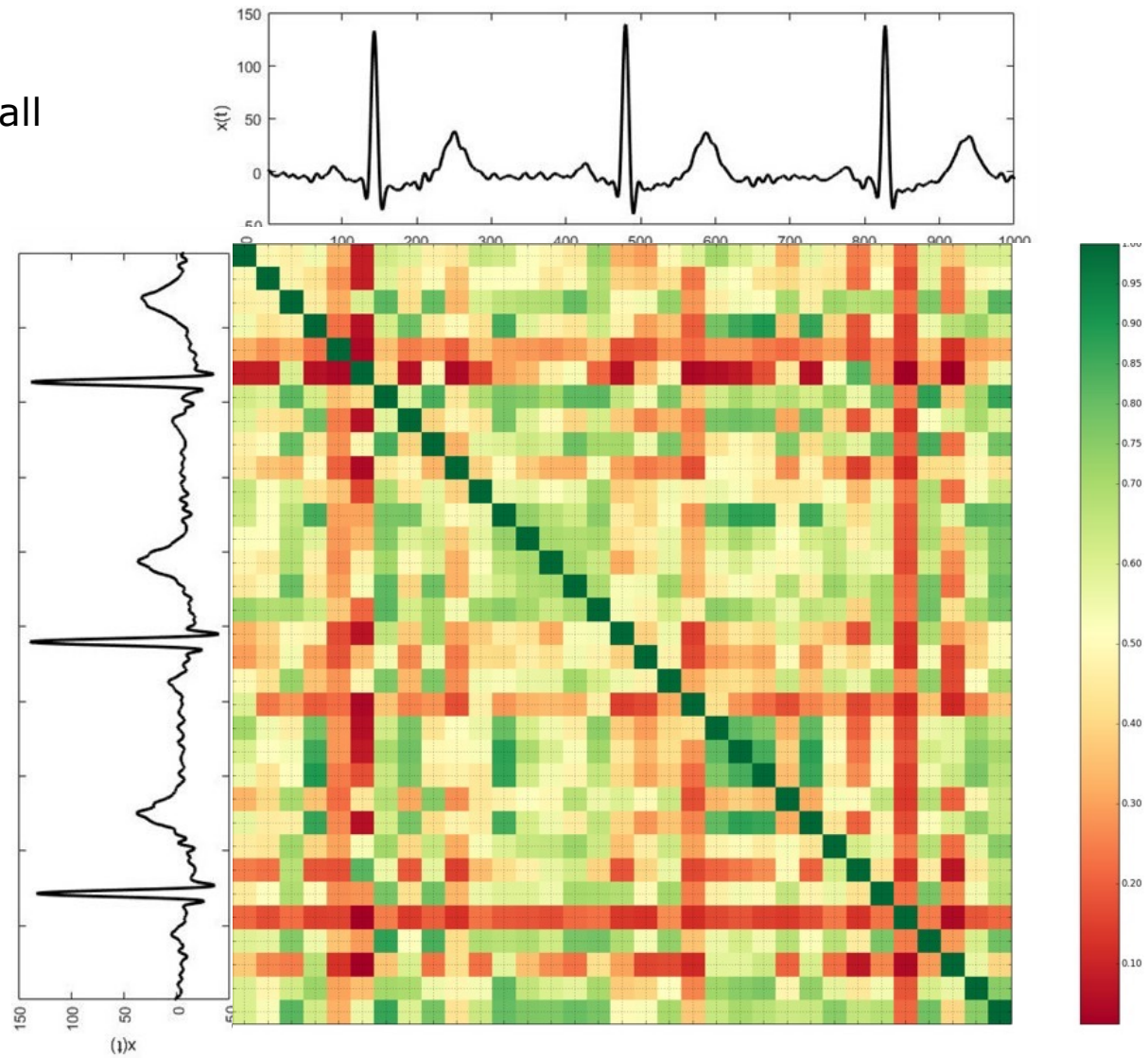
X2	X3	X4	X5
----	----	----	----

Distances

D1,2	...	...	...	...	...	...	...	...	...	...	...
------	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----

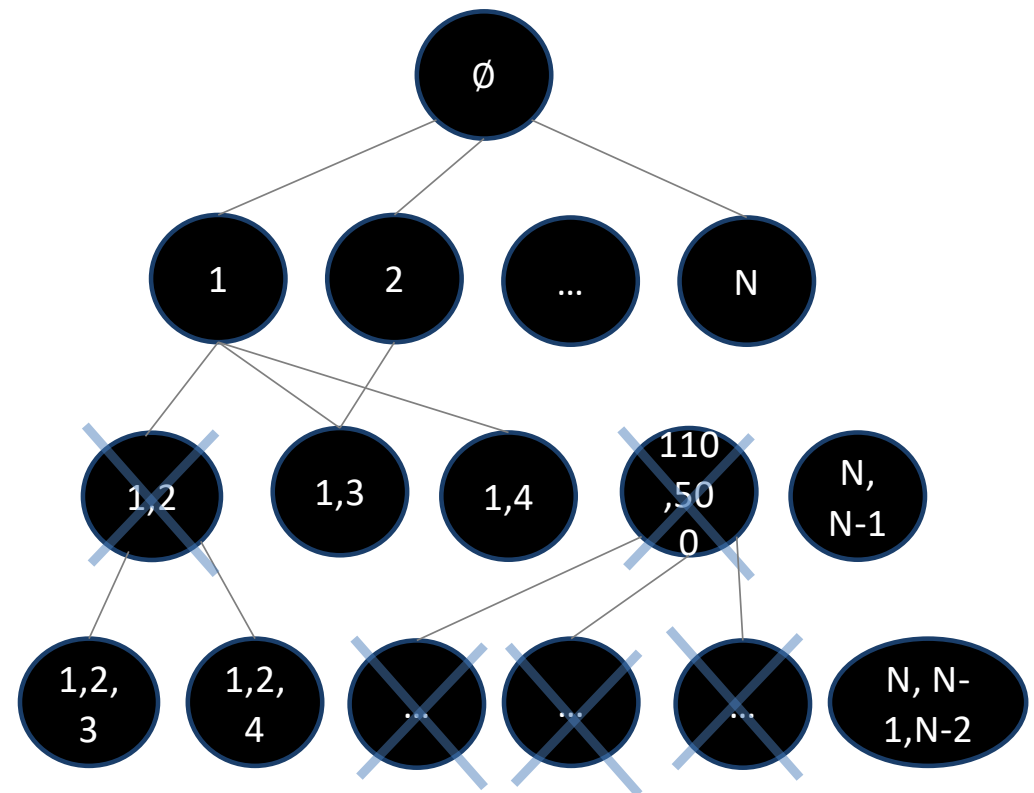
Matrix Profile Distances

D1	...	...	...	...	...	...	...	...	...	...	...
----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----	-----



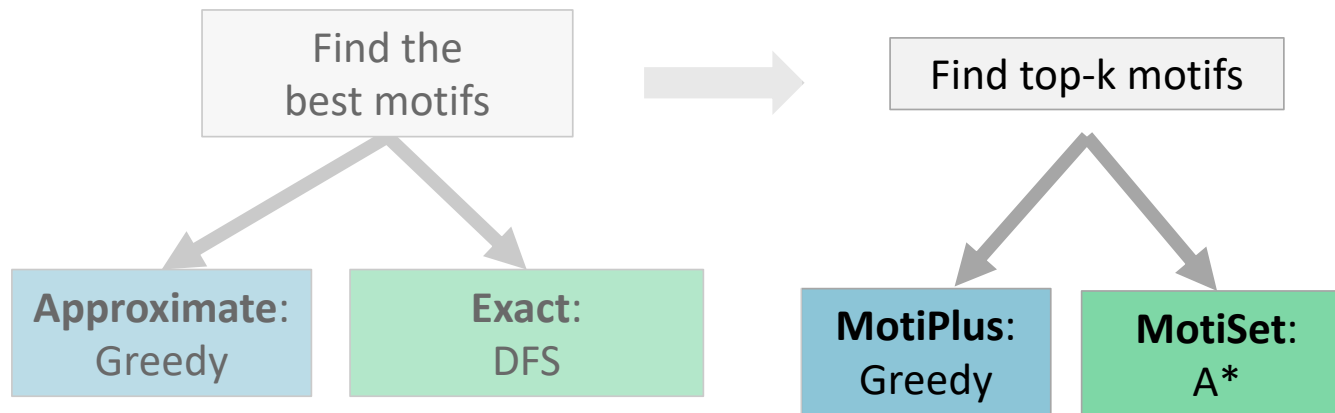
Optionally, **exact search** for motiflets

- **Depth-first-search** to find set of occurrences $\{\mathbf{A}, \mathbf{B}, \mathbf{C}\}$ having the lowest distance
- **Pruning**
 - Initialize **upper bound on distance** using approximate solution
 - **Based** on **monotonicity** of distance
 - **Trivial overlap**
- **Results**
 - **Feasible** up to $K=15$
 - **Greedy search** is highly accurate

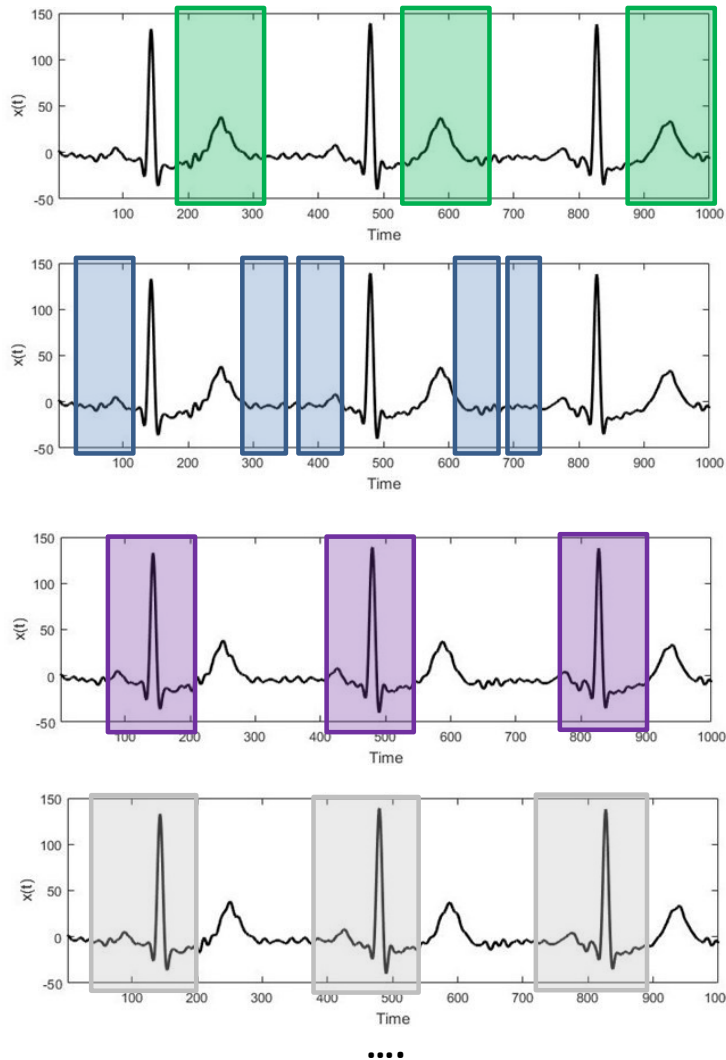


Second contribution

- Novel algorithms for **finding** the **best top-k collection of motifs**, greedy or **anytime**



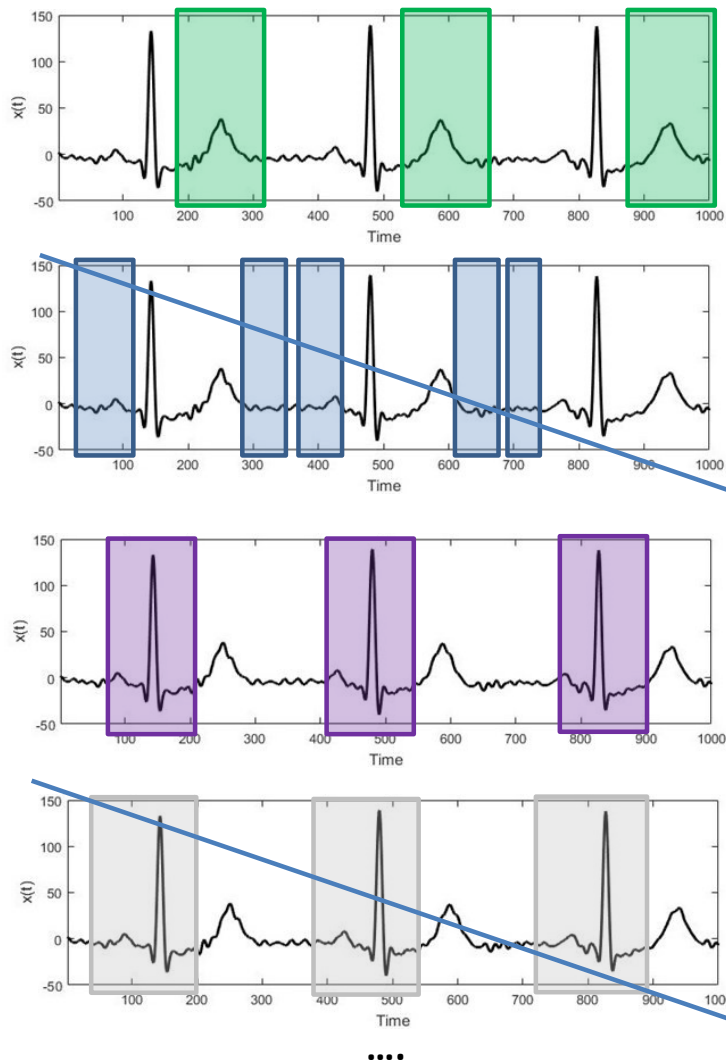
Top-K motifs?



Which 2 do we want to keep?
Many small motifs or few large motifs?
We don't want **redundant** motifs...

Multiple **competing quality objectives**:

- High intra-similarity
- High diversity or low inter-similarity



Multiple **competing quality objectives**:

- High intra-similarity
- High diversity or low inter-similarity

1. **MotiPlus**: Greedy search

- Time complexity: $O(n^2)$

2. **MotiSet**: A^* search

- Time complexity: $O(n^m)$
where m is the total number of motifs

MotiPlus (Greedy Algorithm)

- Find best motif add it, then add next best motif that is not redundant, etc.

- Efficient: m iterations

1. Compute distance matrix and find local motiflets

2. Ranks **local motiflets**:

- lowest extent
- Not overlapping in time domain

3. Apply **filters**:

- **k-closed**
- **self-sufficient**

Algorithm 1: MOTIPLUS: Discovering the top- m motifs

Input: A time series T , motif length l , number of motifs m , the maximum motif size k_{max}

Result: Set S with up to m motifs

```

1 procedure MOTIPLUS( $T, l, m, k_{max}$ )
2    $D \leftarrow \text{CALC\_DISTANCE\_MATRIX}(T, l)$ 
3    $index \leftarrow \text{K\_MOTIFLETS\_HEAP}(T, D, l, k_{max})$ 
4    $S \leftarrow \{\}$ 
5   for  $i \leftarrow 1$  to  $m$  do
6      $S_i \leftarrow \text{BESTMOTIFLET}(T, index, S, k_{max})$ 
7     if  $S_i = \emptyset$  then
8       break
9      $S \leftarrow S \cup \{S_i\}$ 
10  return  $S$ 

11 procedure BESTMOTIFLET( $T, index, S, k_{max}$ )
12   $C \leftarrow \{\}$ 
13  for  $k \leftarrow 2$  to  $k_{max}$  do
14    for  $S_k \in index[k]$  sorted on extent do
15      if  $\text{non-redundant}(S_k, S) \wedge k\text{-closed}(S_k) \wedge \text{self-sufficient}(S_k)$  then
16         $C \leftarrow C \cup \{S_k\}$ 
17        break
18    else
19       $index[k] \leftarrow index[k] \setminus \{S_k\}$ 
20  if  $C = \{\}$  then
21    return  $\emptyset$ 
22   $S \leftarrow \underset{S_k \in C}{\text{argmax}} \text{ elbow}(k)$ 
23  return  $S$ 

```

MotiSet (A* Algorithm)

- Find optimal set of motifs with minimal value

$$\forall S' \subseteq S : f(S) \leq f(S')$$

- Value of set of motifs:

$$f(S) = \text{norm}(\text{smprd}(S)) - \gamma \cdot \text{norm}(\text{smod}(S))$$

– High intra-motif similarity (**compactness**)

$$\text{smprd}(S) = \frac{\sum_{k=1}^m \sum_{i=1}^{|S_k|} \text{min_dist_in}(T_{i,l}, S_k)^2}{\sum_{k=1}^m |S_k|}$$

– Low inter-motif similarity (**diversity**)

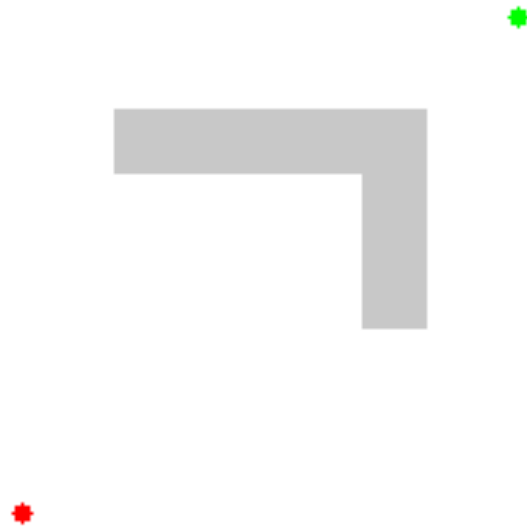
$$\text{smod}(S) = \sum_{k=1}^m \text{min_dist_out}(S_k, T)^2$$

$$\text{min_dist_in}(T_{i,l}, S_k) = \min(\{z\text{-ED}(T_{i,l}, T_{j,l}) \mid T_{j,l} \in S_k \wedge T_{i,l} \neq T_{j,l}\})$$

$$\text{min_dist_out}(S_k, T) = \min(\{z\text{-ED}(T_{i,l}, T_{j,l}) \mid T_{i,l} \in S_k \wedge T_{j,l} \in T\})$$

- In practice we stop the **anytime** search early under a search budget (e.g. 10.000 iterations)

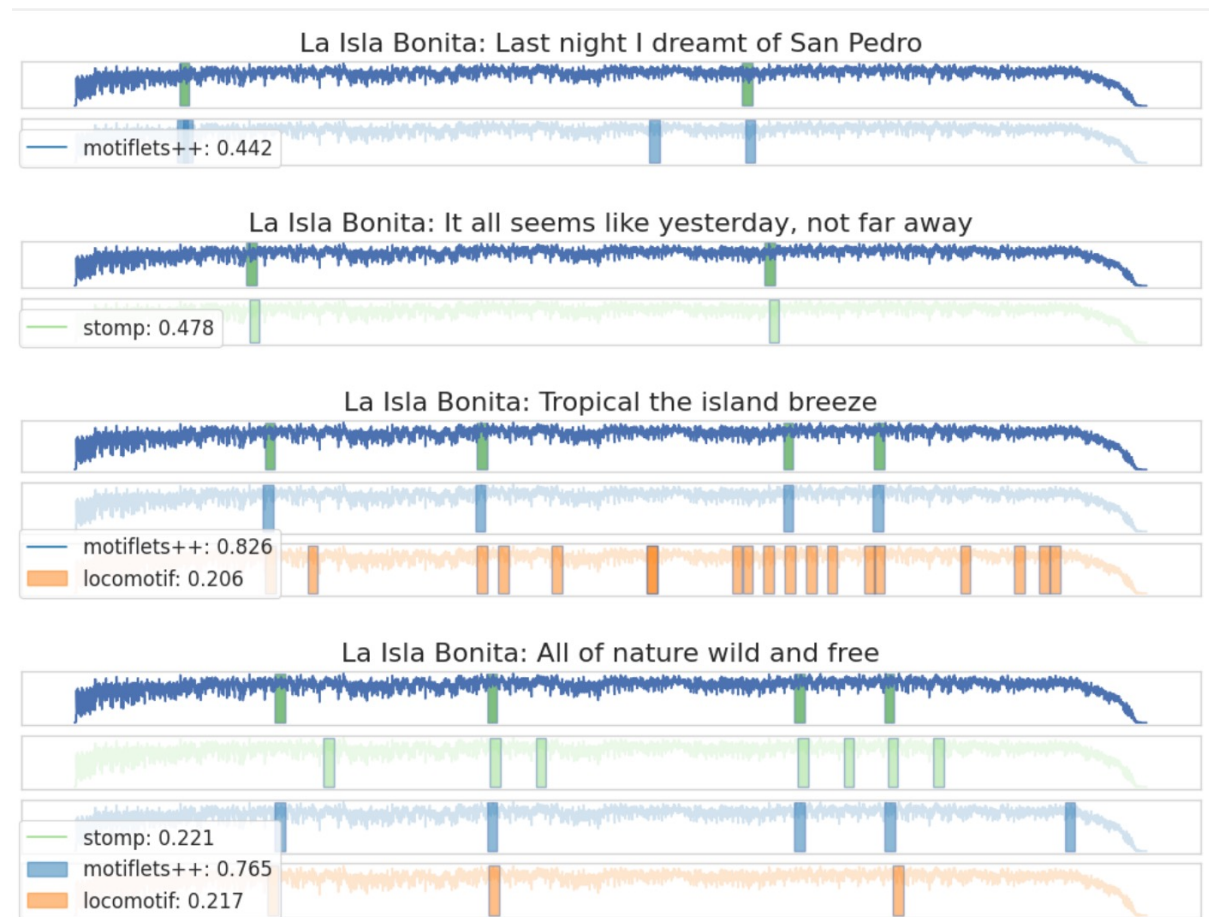
MotiSet (A* Algorithm)



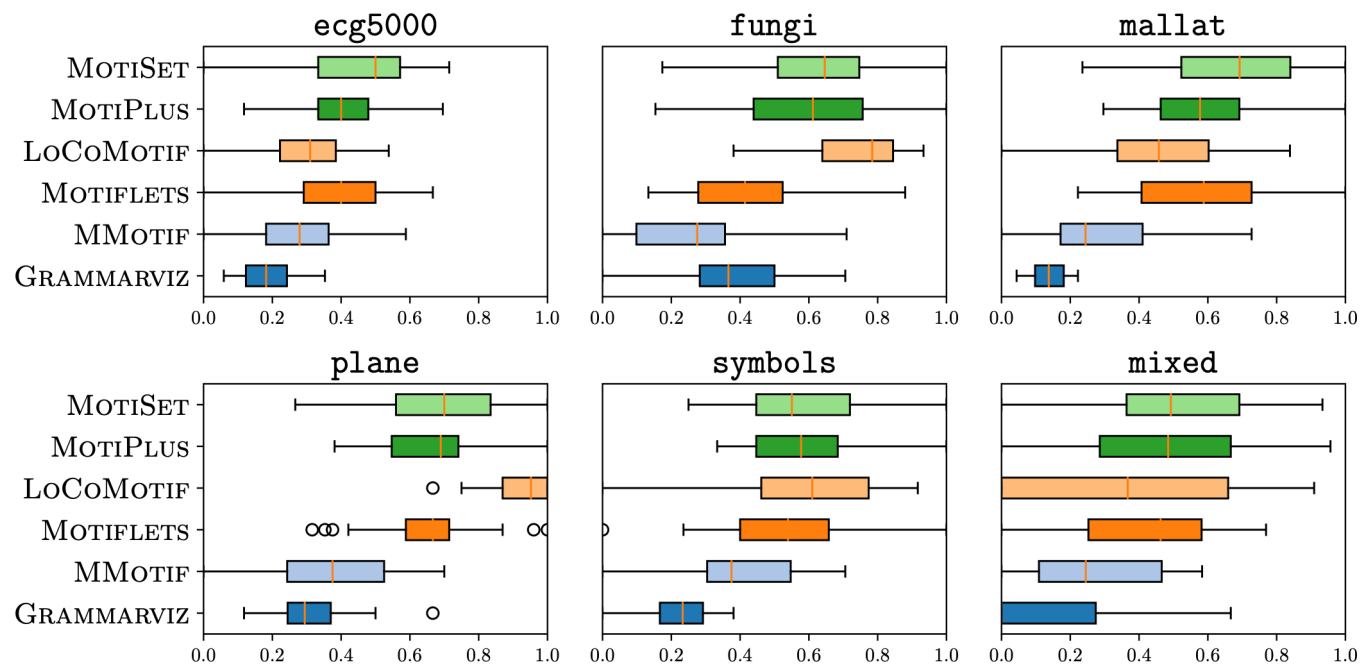
A* search [Hart, 1968]

Results: Music Data

- **MotiPlus** uncovers **recurring** lyrics
- Higher **Jaccard similarity** than state-of-the-art methods
- More **similar** and **diverse motifs**, **better coverage**
- **Example:** La Isla Bonita by Madonna



Results: TSMD-benchmark [Van Wesenbeeck, 2026]



Results: TSMD-benchmark

Dataset	F1	Precision	Recall
GRAMMARVIZ	0.234 (+-0.097)	0.264 (+-0.196)	0.321 (+-0.088)
MMOTIF	0.314 (+-0.058)	0.400 (+-0.106)	0.323 (+-0.061)
MOTIFLETS	0.498 (+-0.104)	0.719 (+-0.125)	0.403 (+-0.094)
LoCoMOTIF	0.565 (+-0.232)	0.604 (+-0.268)	0.564 (+-0.195)
MOTIPLUS	0.552 (+-0.094)	0.543 (+-0.125)	0.621 (+-0.087)
MOTISET	0.586 (+-0.097)	0.773 (+-0.068)	0.531 (+-0.108)

- **MotiSet** has highest F1 and precision. **MotiSet** is slower taking 30 min on average.
- **MotiPlus** has highest recall. Runtime **MotiPlus** is under 2 min.
- **LoCoMotif** highest **F1** on some datasets possibly due to its use of **Dynamic Time Warping**. Takes 10 min on average.

FRM Miner

- To discover motifs in a **database of time series** rather than one series
- Find motifs with varying lengths, varying size and levels of support
- Based on **sequential pattern mining** and post-processing to find highly similar motifs



Motif discovery: key take-aways



Benchmark gains
on TSMD-benchmark



MotiSet:
+27%
precision



MotiPlus:
+10%
recall



Why this matters



- better motif quality



- more useful
top-k results



- clear trade-off
between accuracy
and efficiency



Optimality guarantees



Exact motif search
guarantees
optimal motifs



A* search
guarantees
optimal motif sets



Motif discovery benefits from better quality measures,
efficient search, and exact optimization when needed.

Questions?

Len.feremans@uhasselt.be
Bitbucket.org/len_feremans



FWO grant 12B0V24N