

Data Series Management: Fulfilling the Need for Big Sequence Analytics



Themis Palpanas



Paris Descartes University
Institut Universitaire de France

References

- papers

- ◻ **Data Series Management: Fulfilling the Need for Big Sequence Analytics.** ICDM 2018
 - <http://www.mi.parisdescartes.fr/~themisp/publications/icde18-sms.pdf>
- ◻ **ULISSE: ULtra compact Index for Variable-Length Similarity Search in Data Series.** ICDE 2018
 - <http://www.mi.parisdescartes.fr/~themisp/publications/icde18-ulisse.pdf>
- ◻ **DPiSAX: Massively Distributed Partitioned iSAX.** ICDM 2017
 - <http://www.mi.parisdescartes.fr/~themisp/publications/icdm17-dpisax.pdf>
- ◻ **ADS: The Adaptive Data Series Index.** VLDBJ 2016
 - <http://www.mi.parisdescartes.fr/~themisp/publications/vldb16-ads.pdf>
- ◻ **Big Sequence Management: A Glimpse on the Past, the Present, and the Future.** LNCS, 2016
 - <http://www.mi.parisdescartes.fr/~themisp/publications/sofsem16-bisem.pdf>
- ◻ **Query Workloads for Data-Series Indexes.** KDD 2015
 - <http://www.mi.parisdescartes.fr/~themisp/publications/kdd15-bends.pdf>
- ◻ **RINSE: Interactive Data Series Exploration.** VLDB 2015
 - <http://www.mi.parisdescartes.fr/~themisp/publications/vldb15-rinse.pdf>
- ◻ **Indexing for Interactive Exploration of Big Data Series.** SIGMOD 2014
 - <http://www.mi.parisdescartes.fr/~themisp/publications/sigmod14-ads.pdf>
- ◻ **Beyond One Billion Time Series: Indexing and Mining Very Large Time Series Collections with iSAX2+.** KAIS 2014
 - <http://www.mi.parisdescartes.fr/~themisp/publications/kais14-isax2plus.pdf>
- ◻ **iSAX 2.0: Indexing and Mining One Billion Time Series.** ICDM 2010
 - <http://www.mi.parisdescartes.fr/~themisp/publications/icdm10-billiontimeseries.pdf>

- code and datasets

- ◻ <https://github.com/zoumpatianos/ADS>
- ◻ <http://www.mi.parisdescartes.fr/~themisp/isax2plus/>

- data series toolbox

- ◻ <https://github.com/zoumpatianos/DStat>

- demo

- ◻ <http://daslab.seas.harvard.edu/rinse/>

Acknowledgements

Paris Descartes University

- Michele Linardi
- Anna Gogolou
- Botao Peng
- Karima Echihabi



University of Trento

- Alessandro Camerra

Harvard University

- Stratos Idreos
- Kostas Zoumpatianos

Cornell University/Microsoft

- Yin Lou
- Johannes Gehrke



Inria

- Djamel-Edine Yagoubi
- Reza Akbarinia
- Florent Masseglia

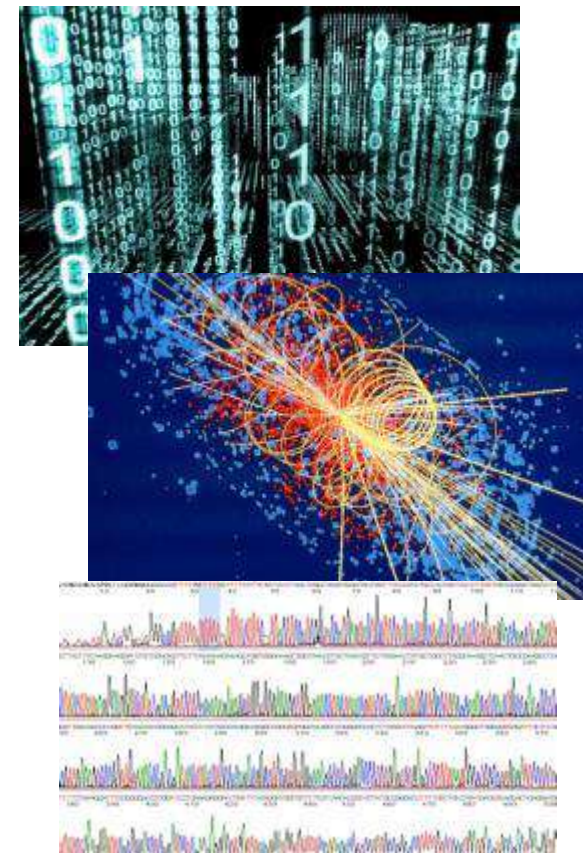
University of California at Riverside

- Jin Shieh
- Eamon Keogh



Executive Summary

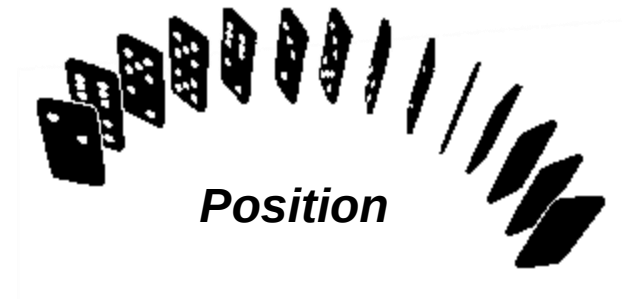
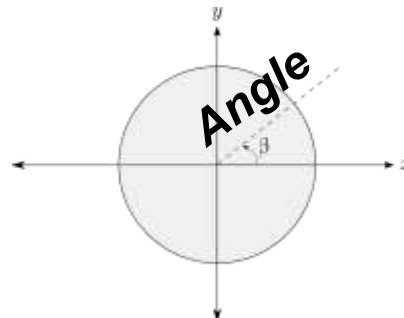
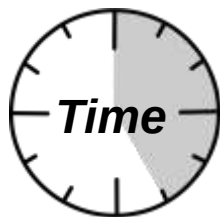
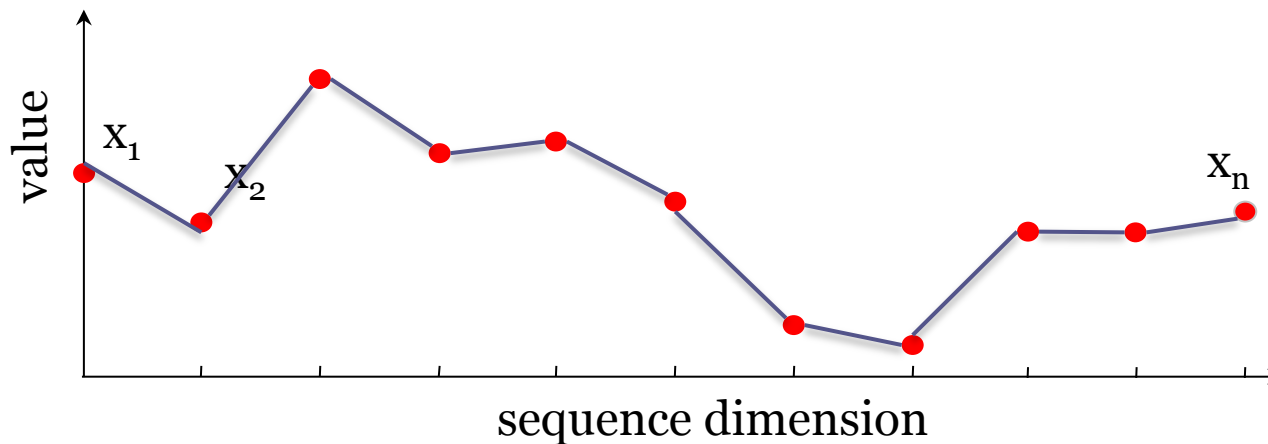
- data collected at unprecedented rates
- they enable data-driven scientific discovery
- lots of these data are sequences
 - takes **days-weeks** to analyze big sequence collections



goal: **analyze big sequences** in **minutes/seconds**

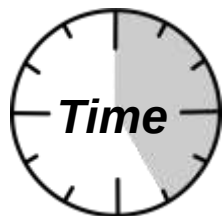
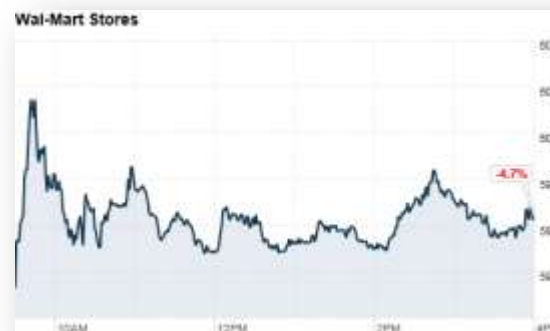
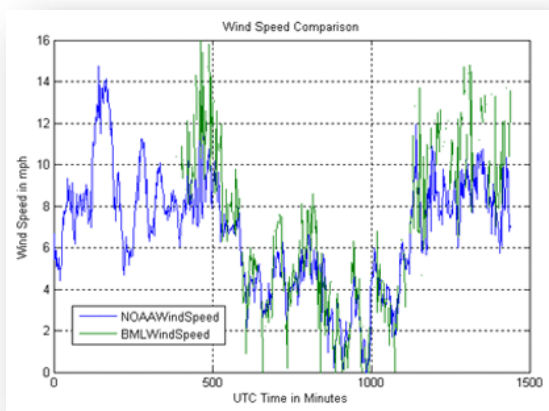
Data series

- Sequence of points ordered along some dimension



Scientific Monitoring

- meteorology, oceanography, astronomy, finance, sociology, ...



Wind speed

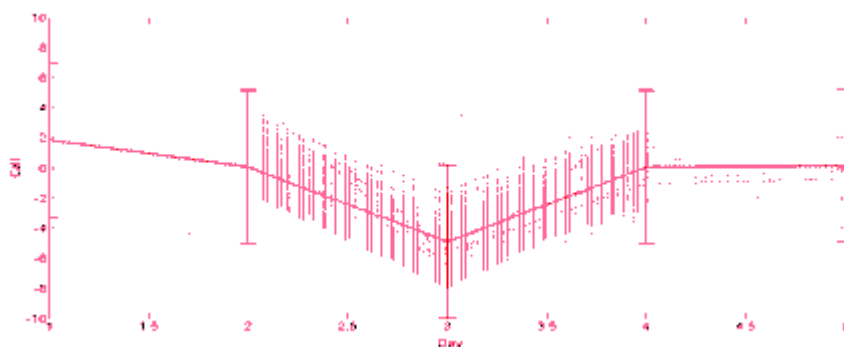
From ocean observing node project
<http://bml.ucdavis.edu/boon/wind.html>

Historical stock quotes

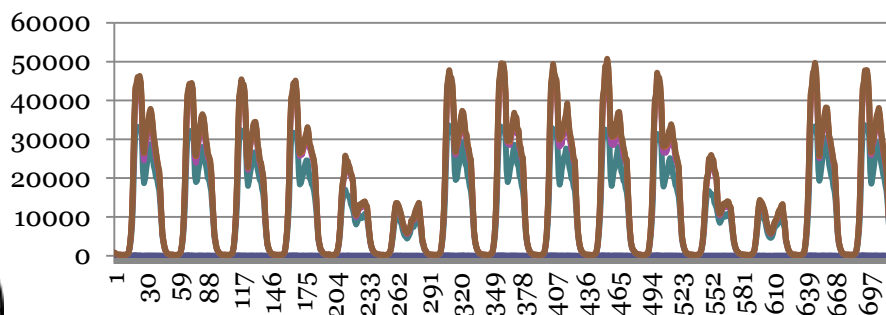
http://money.cnn.com/2012/04/23/markets/walmart_stock/index.htm

Telecommunications

- analysis of **call activity** patterns
 - Telecom Italia



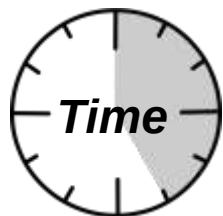
call activity for Easter Monday



average number of calls for 5 smallest clusters

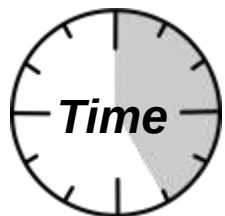
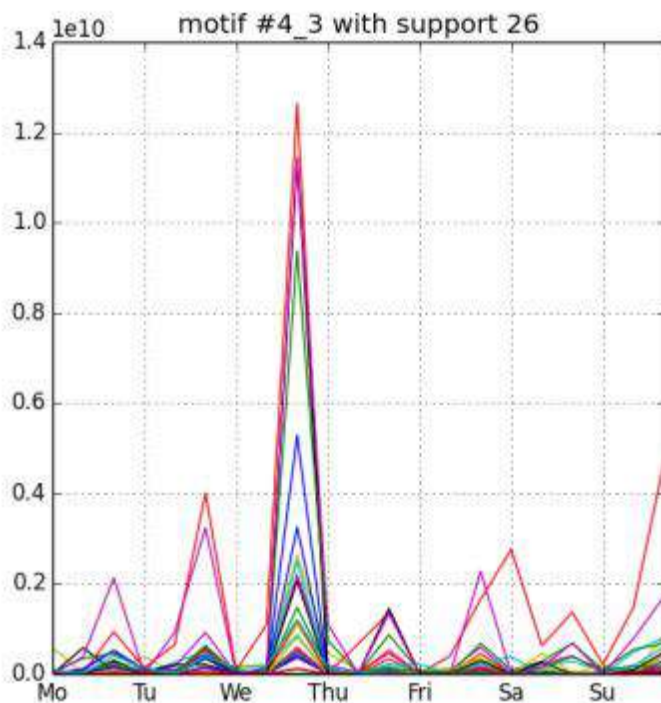


clustermmap of incoming calls time series

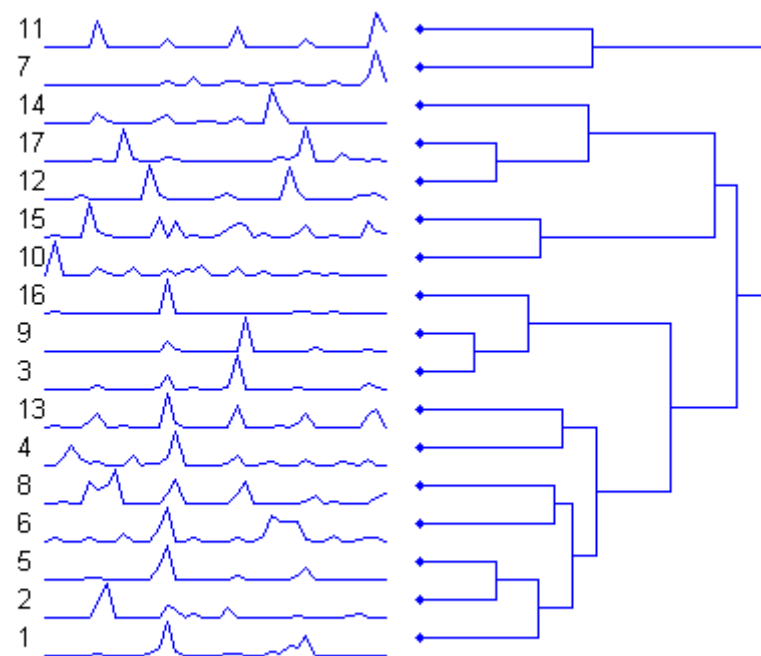


Home Networks

- temporal **usage behavior analysis** of home networks
 - Portugal Telecom



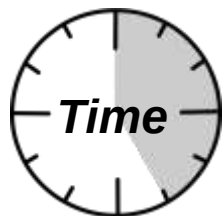
(previously unknown) frequent behavior pattern



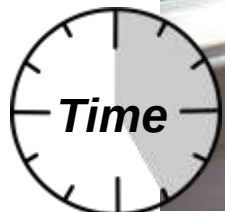
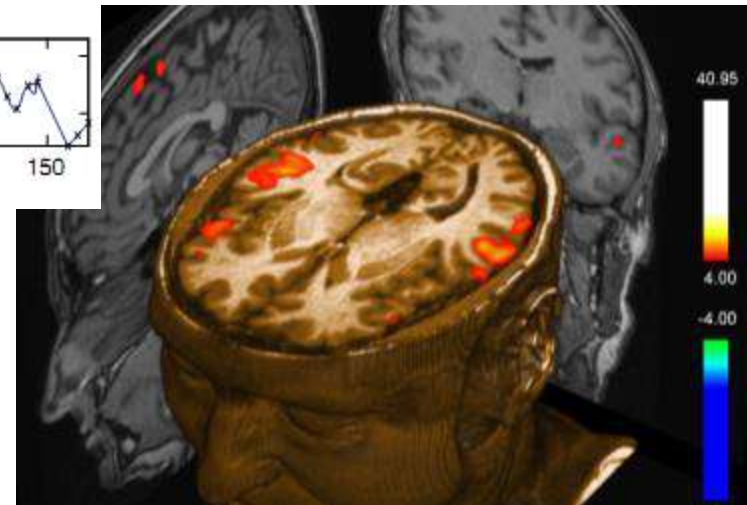
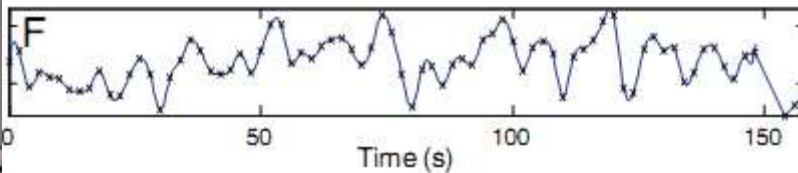
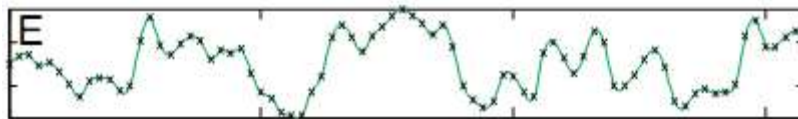
clustering based on user activity patterns

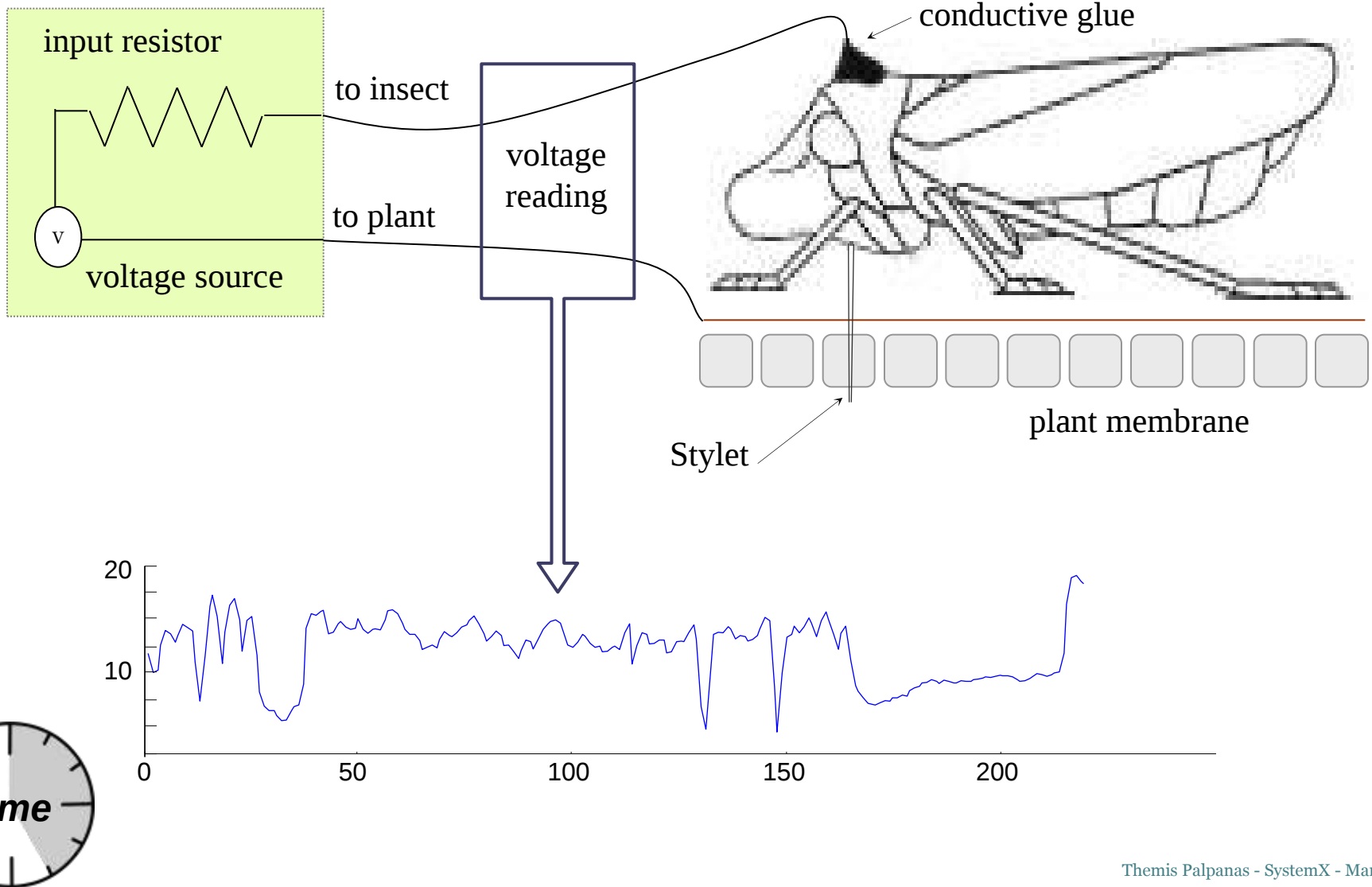
Data Centers

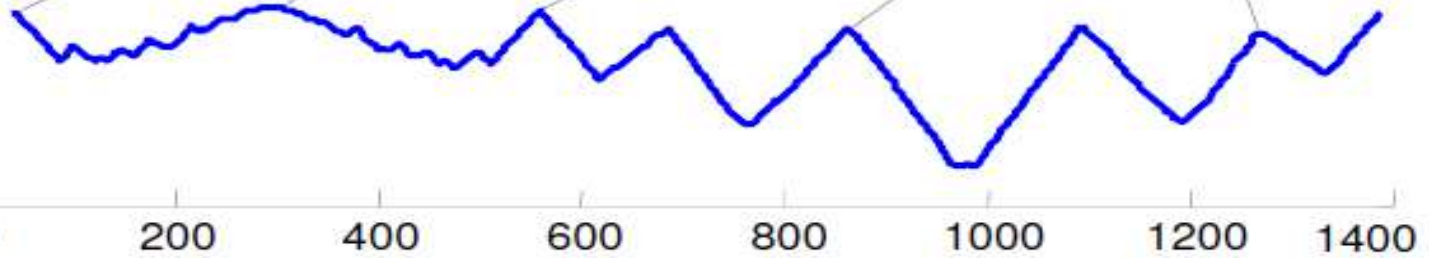
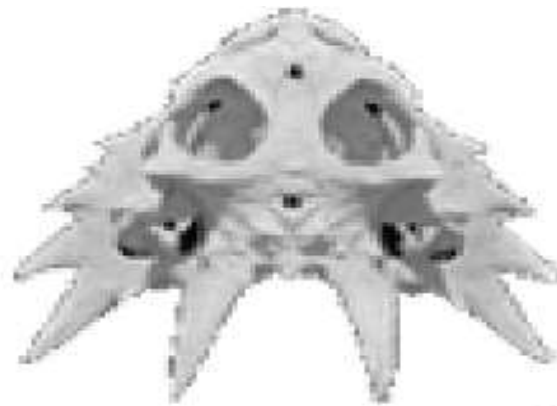
- cloud utilization/operation/health monitoring



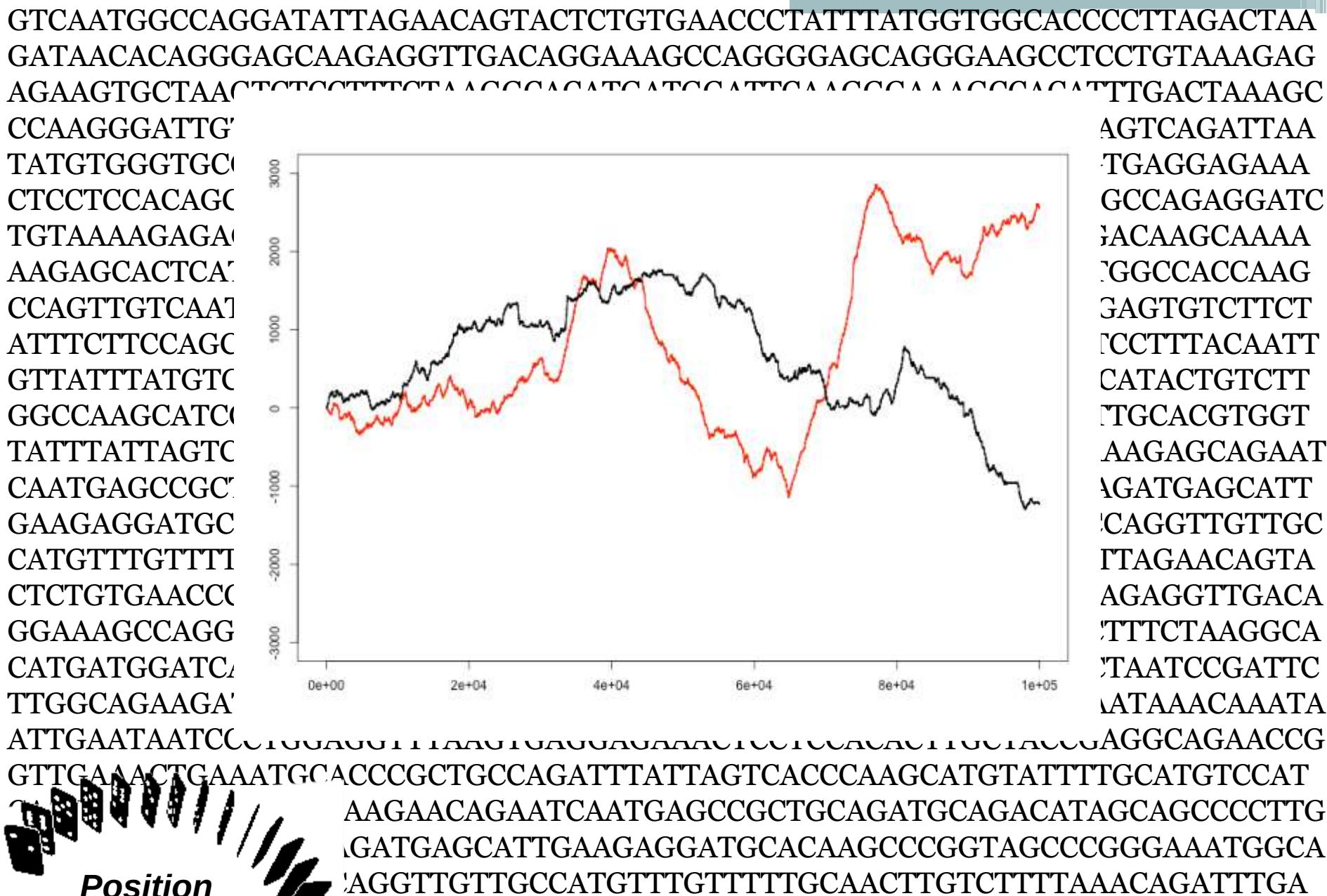
- functional Resonance Magnetic Imaging (fMRI) data
 - primary experimental tool of neuroscientists
 - reveal how different parts of brain respond to stimuli







Angle



What do we want to do with them?

- simple query answering

**select values
in time
interval**

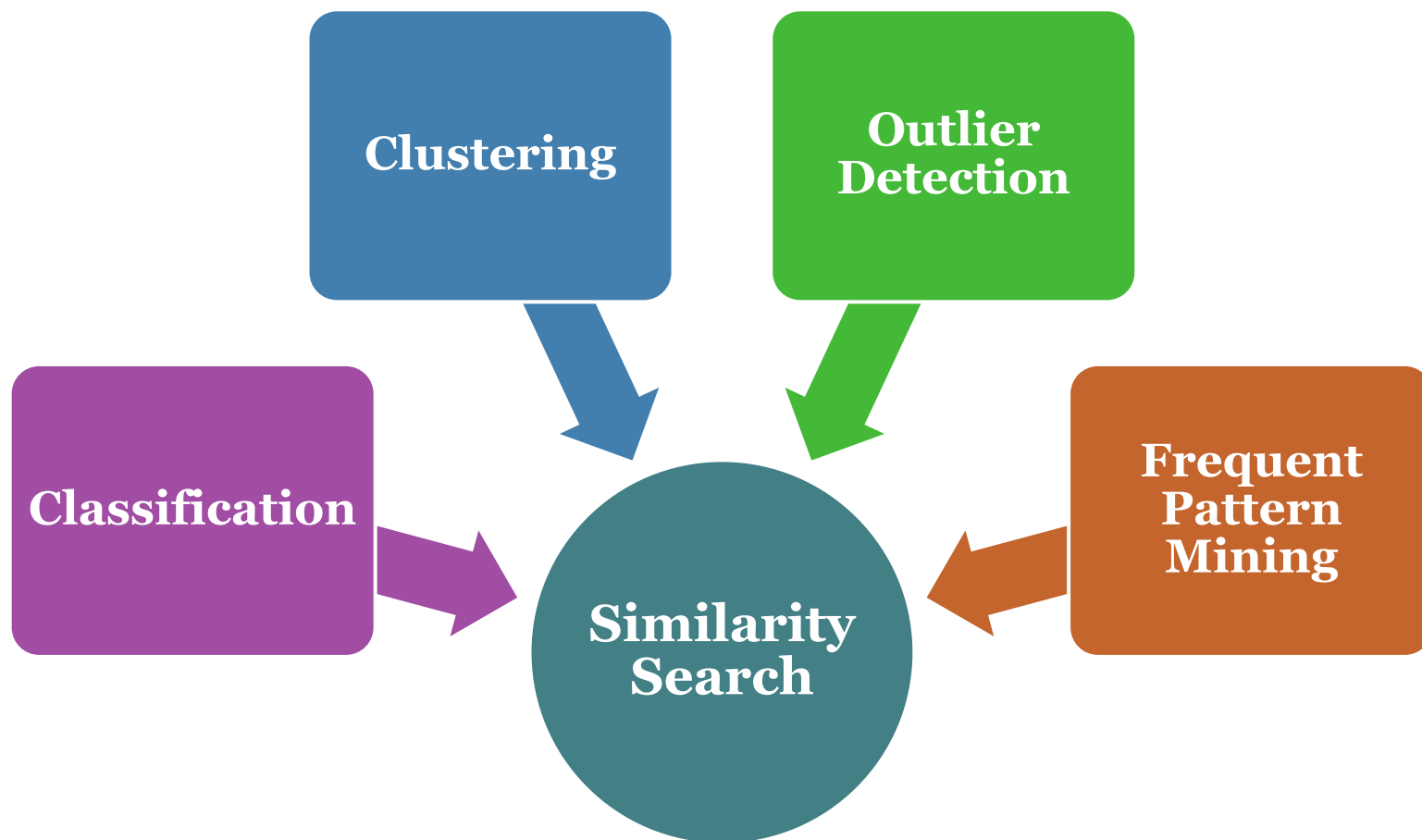
**select values
in some
range**

**select some
data series**

**combinations
of those**

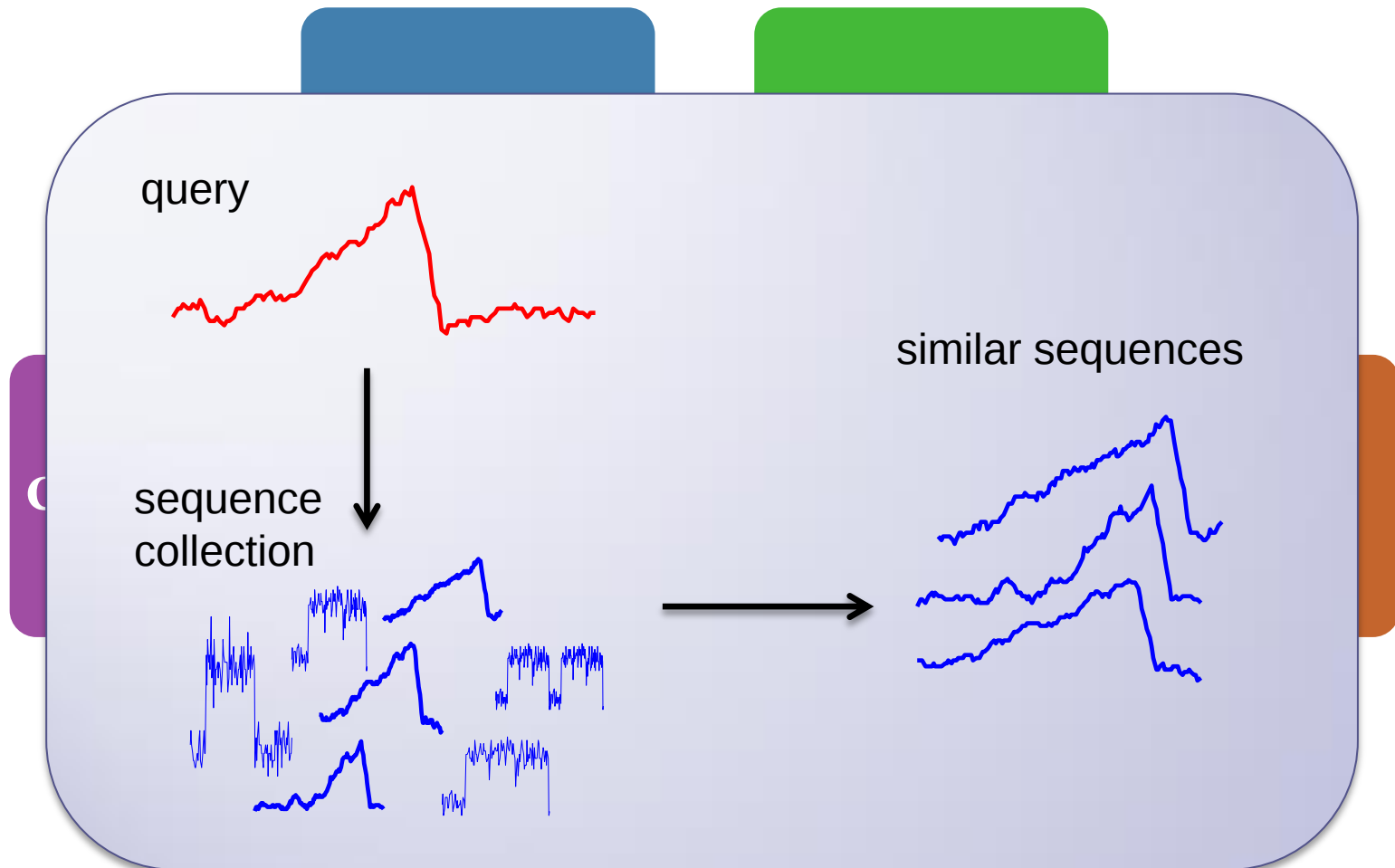
What do we want to do with them?

- complex analytics



What do we want to do with them?

- complex analytics



What do we want to do with them?

- complex analytics

Euclidean

$$D(X, Y) \equiv \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$$

Dynamic Time
Warping (DTW)

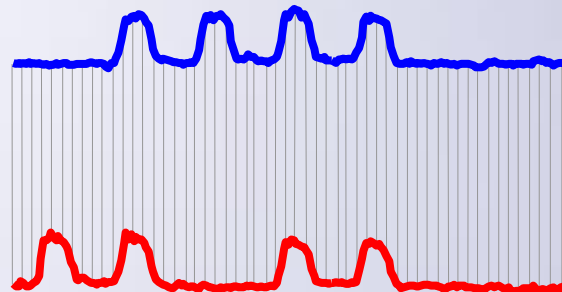
$$D_{dtw}(X, Y) = f(n, m)$$

$$f(i, j) = \|x_i - y_j\| + \min \begin{cases} f(i, j-1) \\ f(i-1, j) \\ f(i-1, j-1) \end{cases}$$

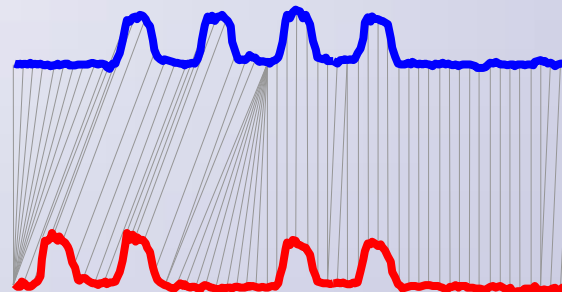
What do we want to do with them?

- complex analytics

Euclidean

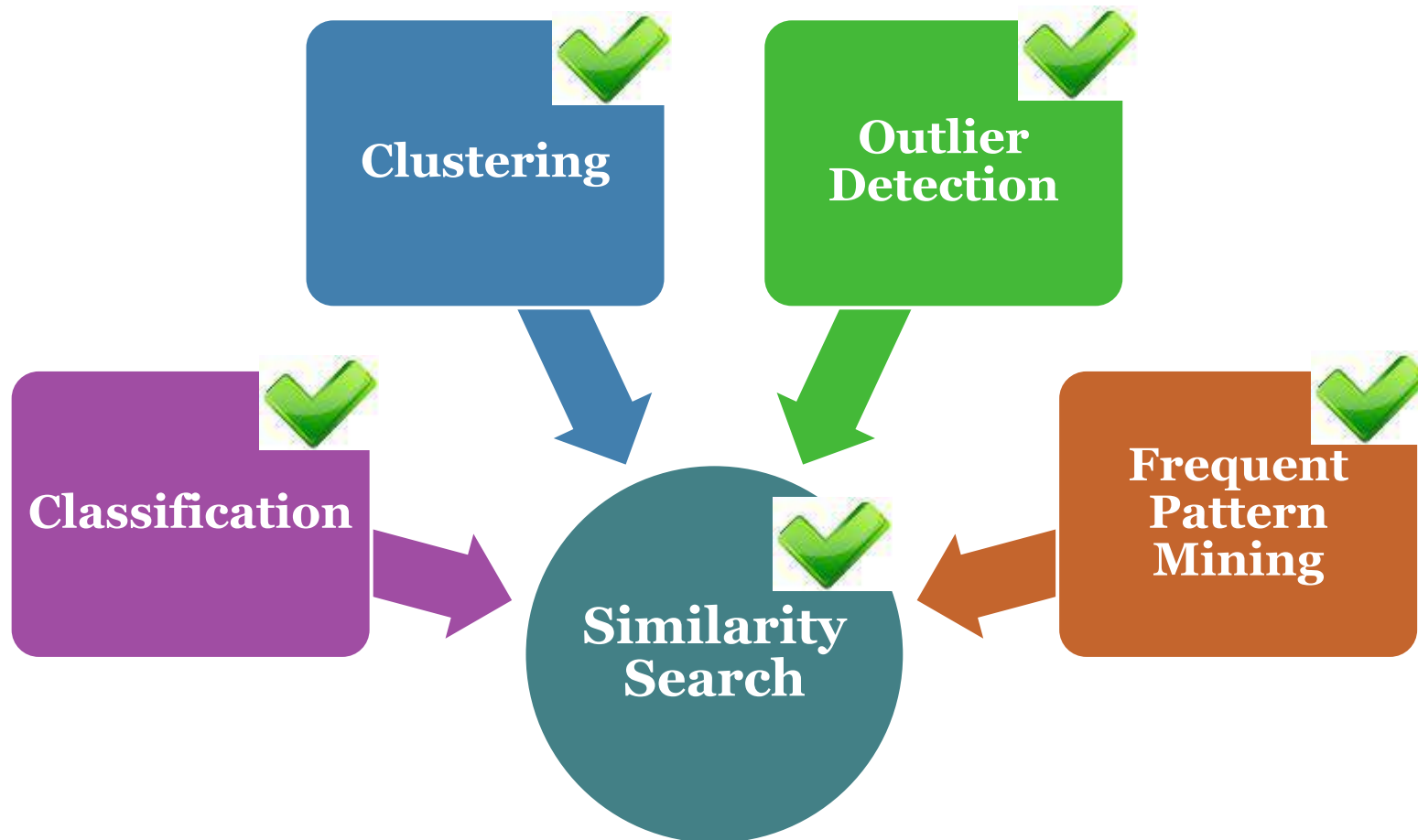


Dynamic Time
Warping (DTW)



What do we want to do with them?

- complex analytics



What do we want to do with them?

- complex analytics



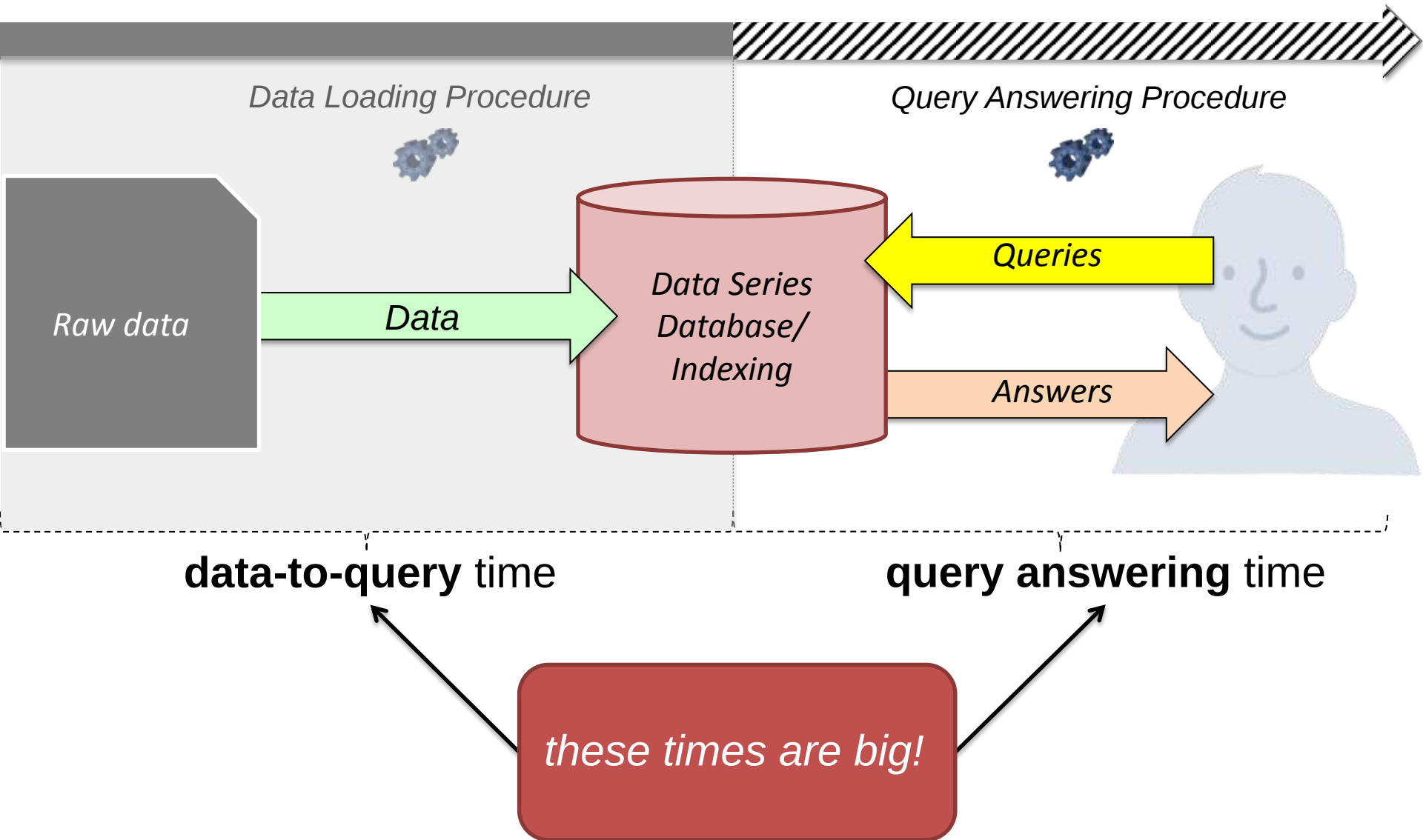
Clustering

Outlier
Detection

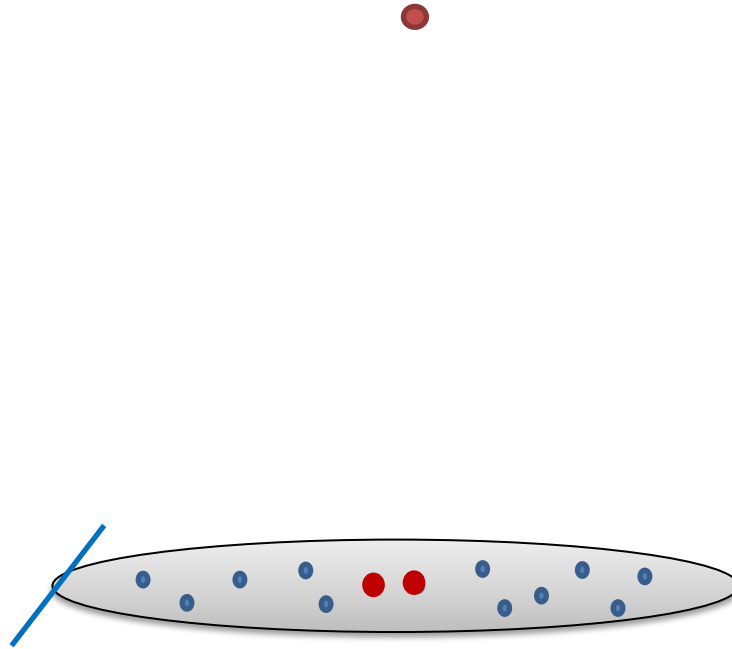
HARD, because of **very high dimensionality:
each data series has 100s-1000s of points!**

even HARDER, because of **very large size:
millions to billions of data series (multi-TBs)!**

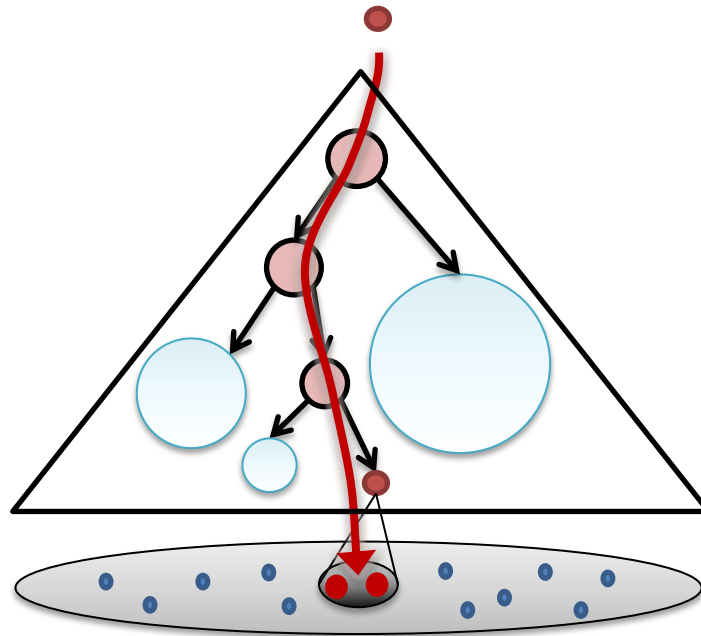
Query answering process



Similarity Search via Serial Scan



Similarity Search via Indexing



Traditional Approaches

answer nearest neighbor queries on a 1TB dataset:

serial scan takes **45 minutes/query**



but building the index takes **too long!**

indexing a 1TB dataset takes **days**

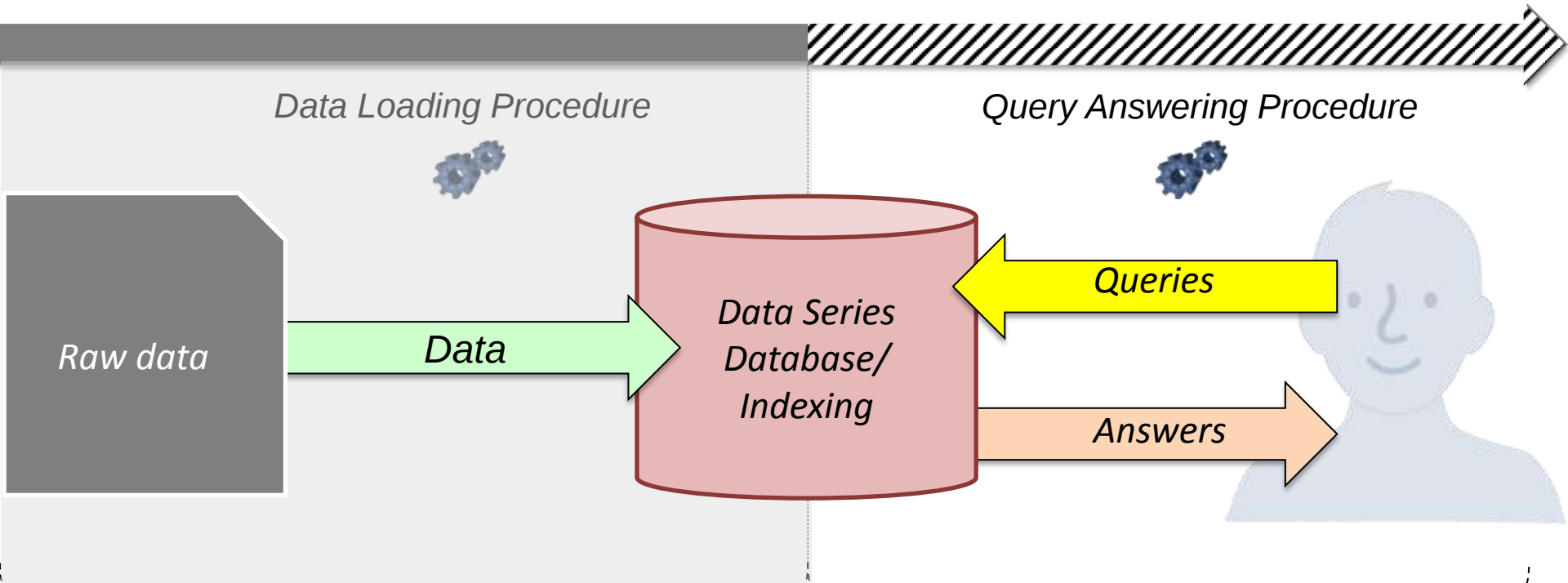


a data series index can
reduce querying time

Query Answering

complex analytics in days...

Query answering process



data-to-query time

query answering time

*we have proposed the
state-of-the-art
solutions for both problems!*

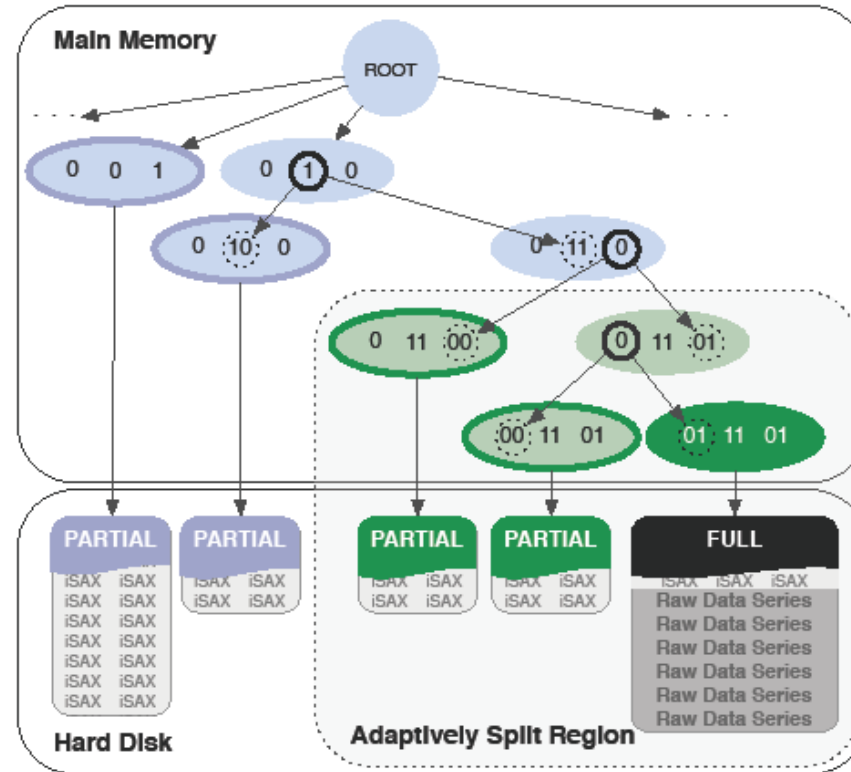
State of the Art Approach: ADS+

Publications

SIGMOD'14

PVLDB'15

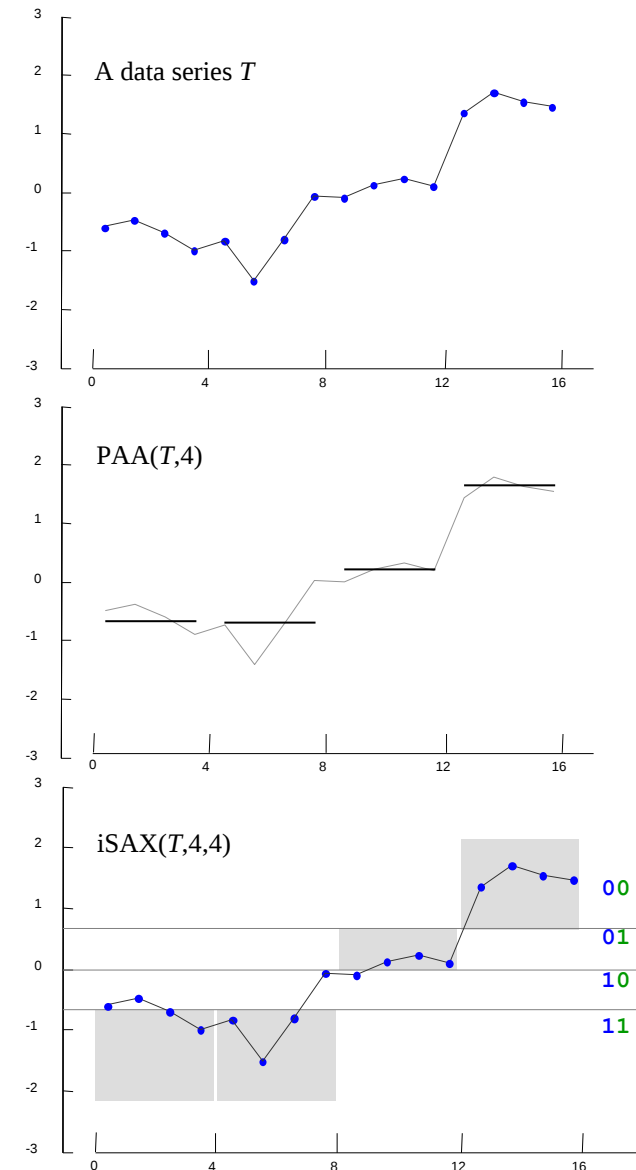
VLDBJ'16



complex analytics in hours!

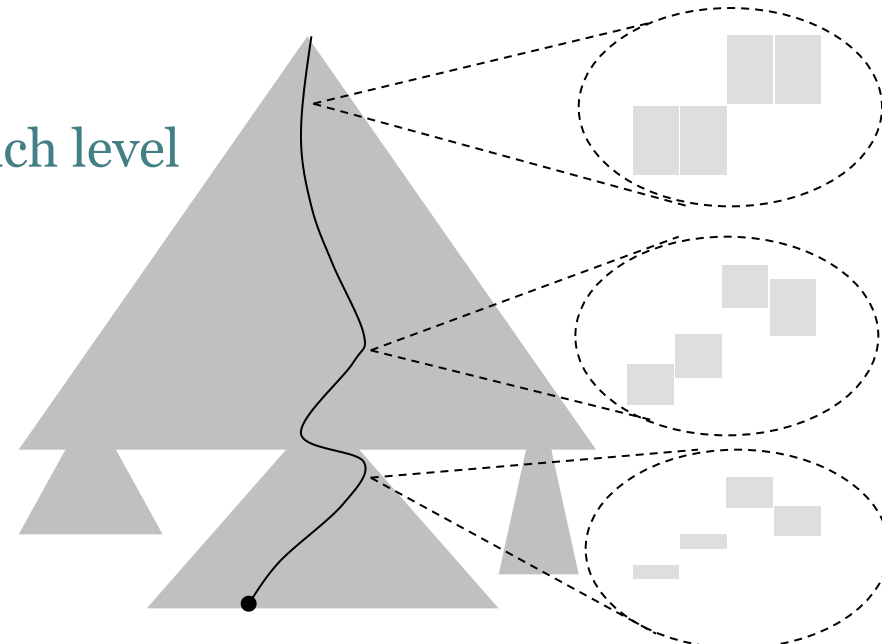
SAX Representation

- **Symbolic Aggregate approXimation (SAX)**
 - **(1)** Represent data series T of length n with w segments using Piecewise Aggregate Approximation (PAA)
 - T typically normalized to $\mu = 0, \sigma = 1$
 - $PAA(T, w) = \bar{T} = \bar{t}_1, \dots, \bar{t}_w$
 - where $\bar{t}_i = \frac{w}{n} \sum_{j=\frac{n}{w}(i-1)+1}^{\frac{n}{w}i} T_j$
 - **(2)** Discretize into a vector of symbols
 - Breakpoints map to small alphabet α of symbols



iSAX Index

- non-balanced tree-based index with non-overlapping regions, and controlled fan-out rate
 - base cardinality ***b*** (optional), segments ***w***, threshold ***th***
 - hierarchically subdivides SAX space until num. entries $\leq th$
- Approximate Search
 - Match iSAX representation at each level
- Exact Search
 - Leverage approximate search
 - Prune search space
 - Lower bounding distance



Drawback of iSAX2+

- cannot start answering queries until *entire* index is built!

Adaptive Data Series Index: ADS+

- **novel paradigm** for building a data series index
 - do not build entire index and then answer queries
 - start answering queries by building the part of the index needed by those queries
- still guarantee **correct answers**

Adaptive Data Series Index: ADS+

Publications

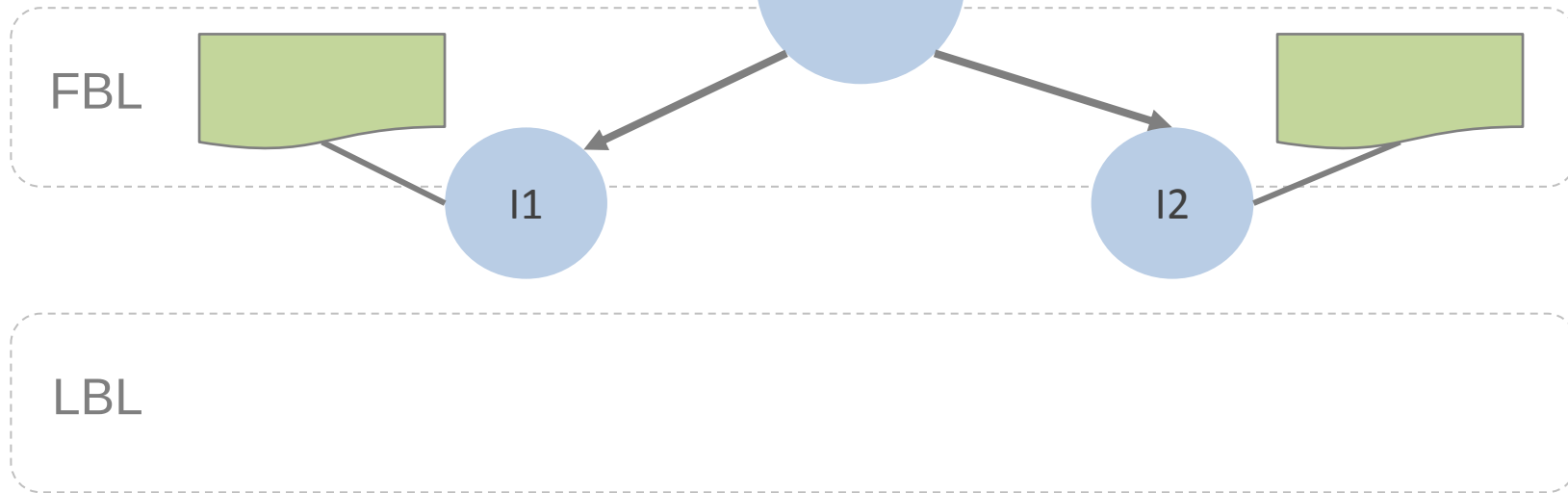
SIGMOD'14

PVLDB'15

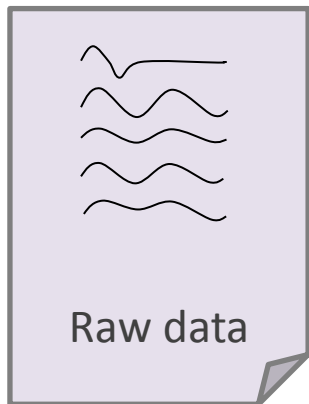
VLDBJ'16

- intuition for proposed solution
 - build the iSAX index using the iSAX representations
 - just like iSAX2+
 - but start with a large leaf size
 - minimize initial cost
 - postpone leaf materialization to query time
 - only materialize (at query time) leaves needed by queries
 - parts that are queried more are refined more
 - use smaller leaf sizes (reduced leaf materialization and query answering costs)

Start building an index with only the iSAX representations

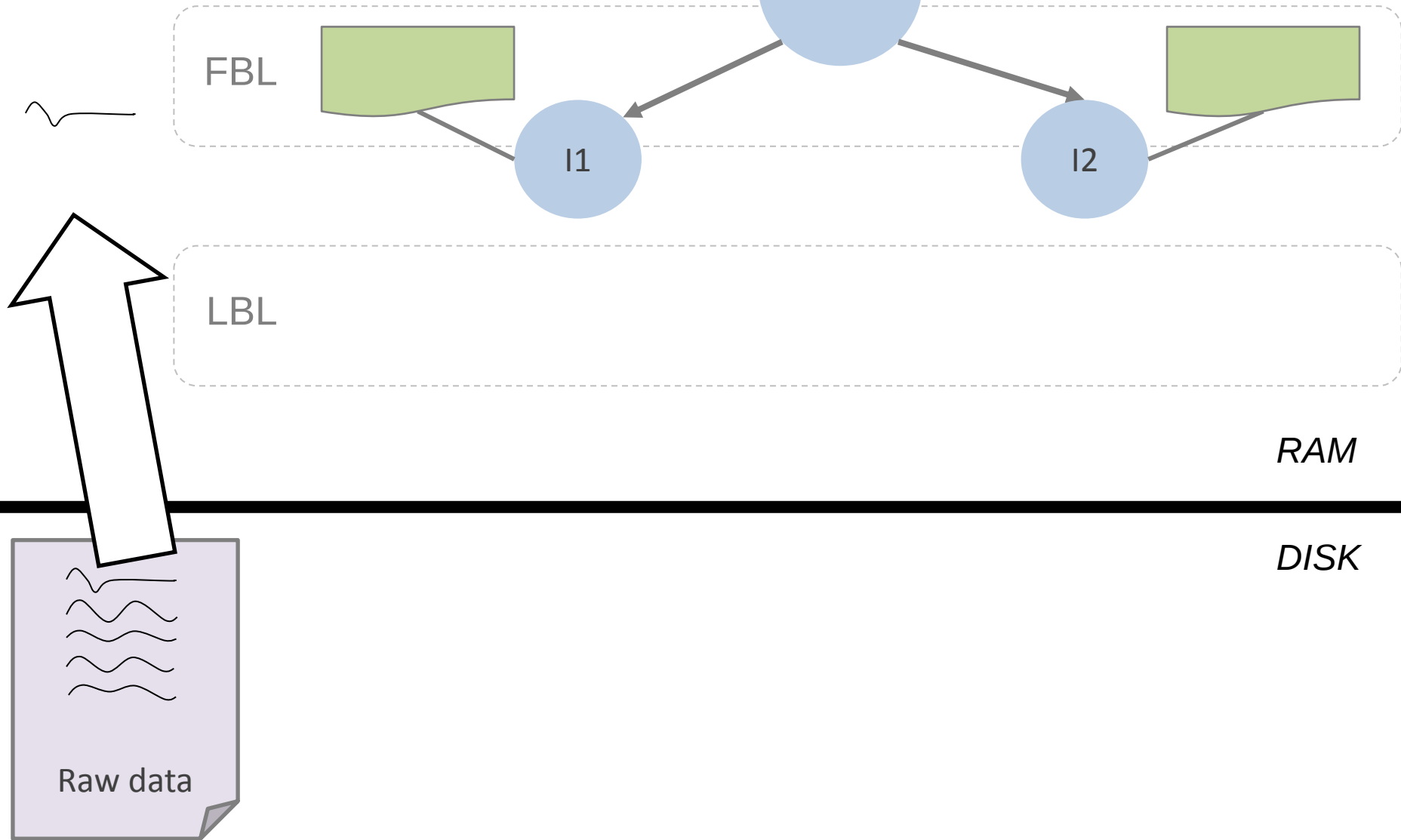


RAM

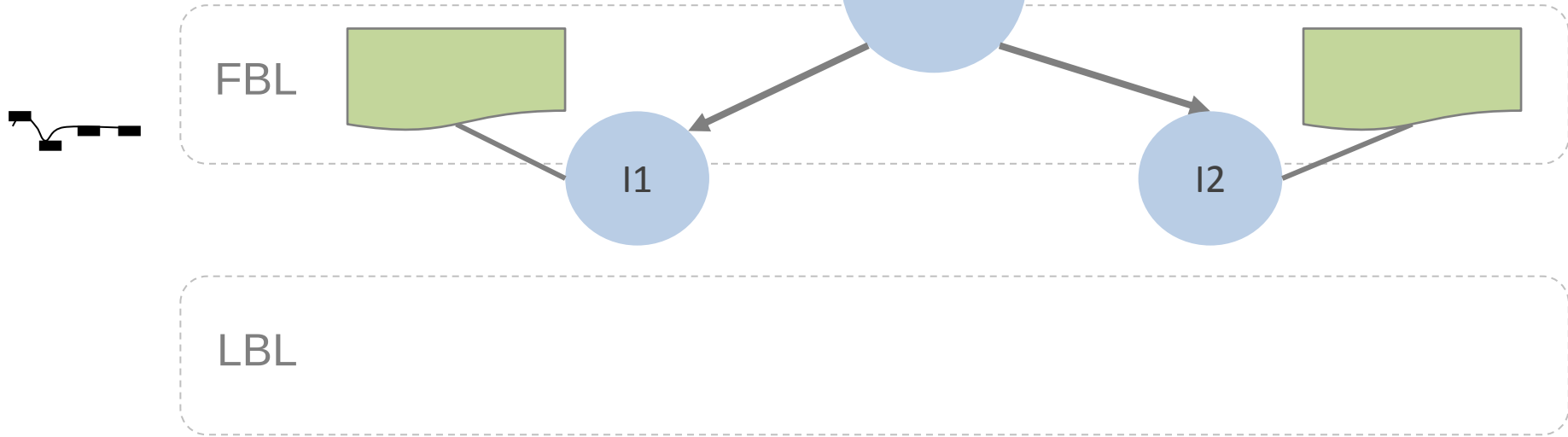


DISK

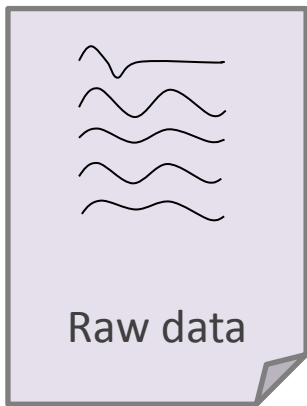
Read the data-series one by one from the raw file



Convert them to iSAX

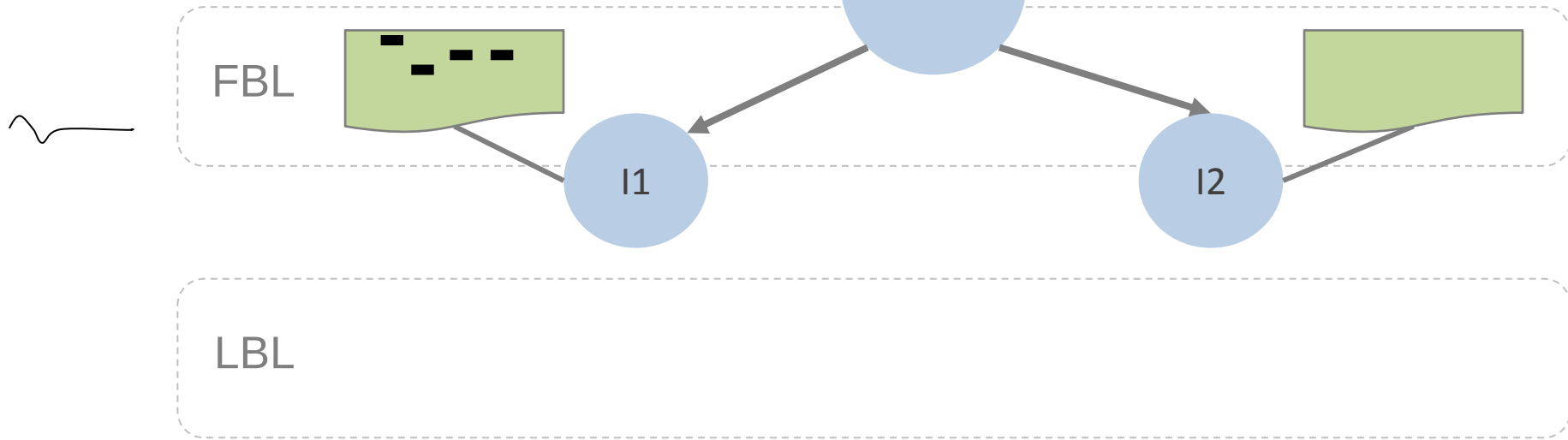


RAM

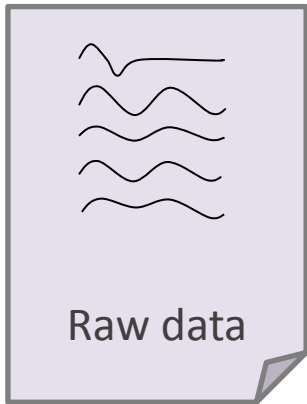


DISK

Store only iSAX in memory (64 times smaller) ~1%

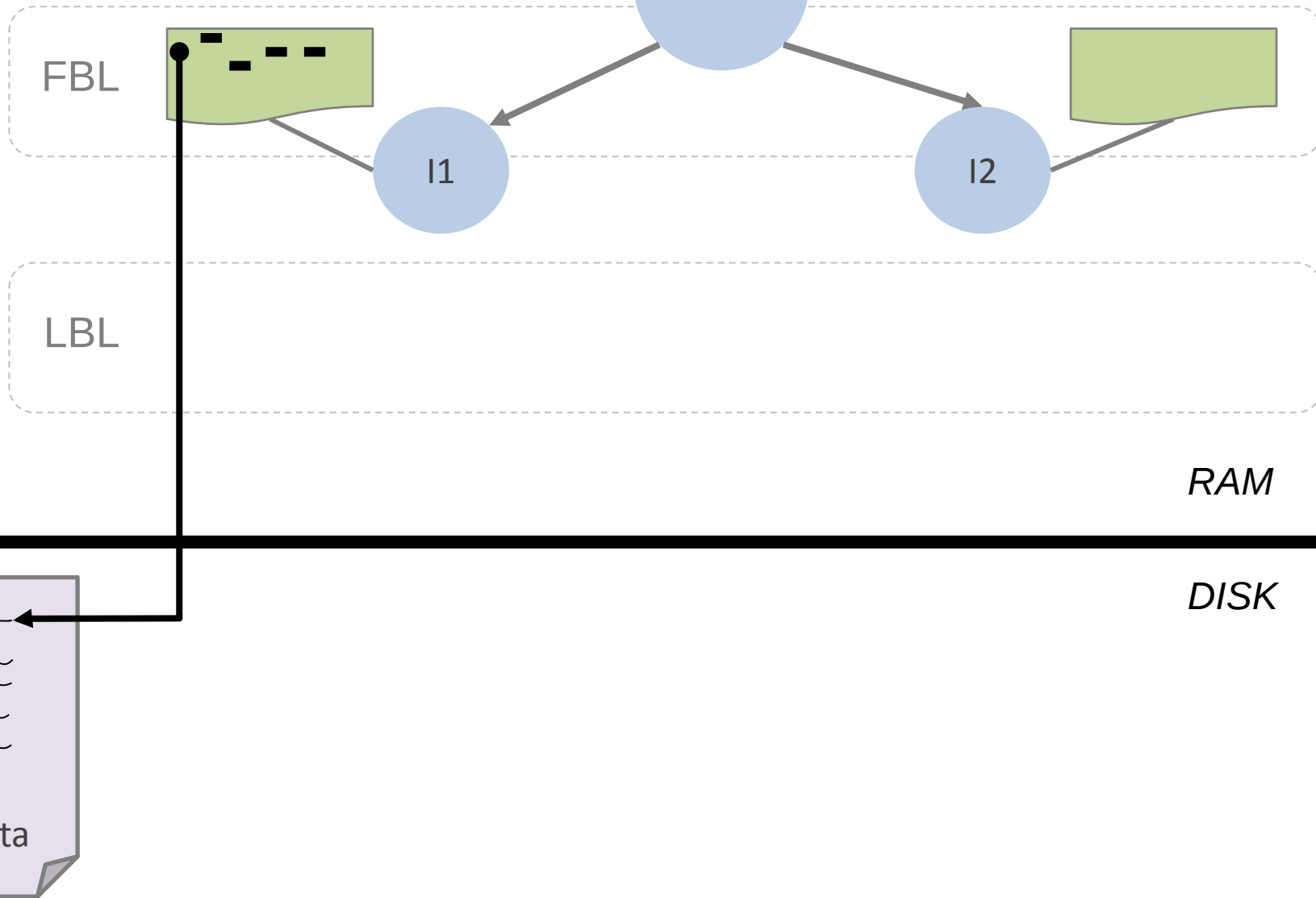


RAM



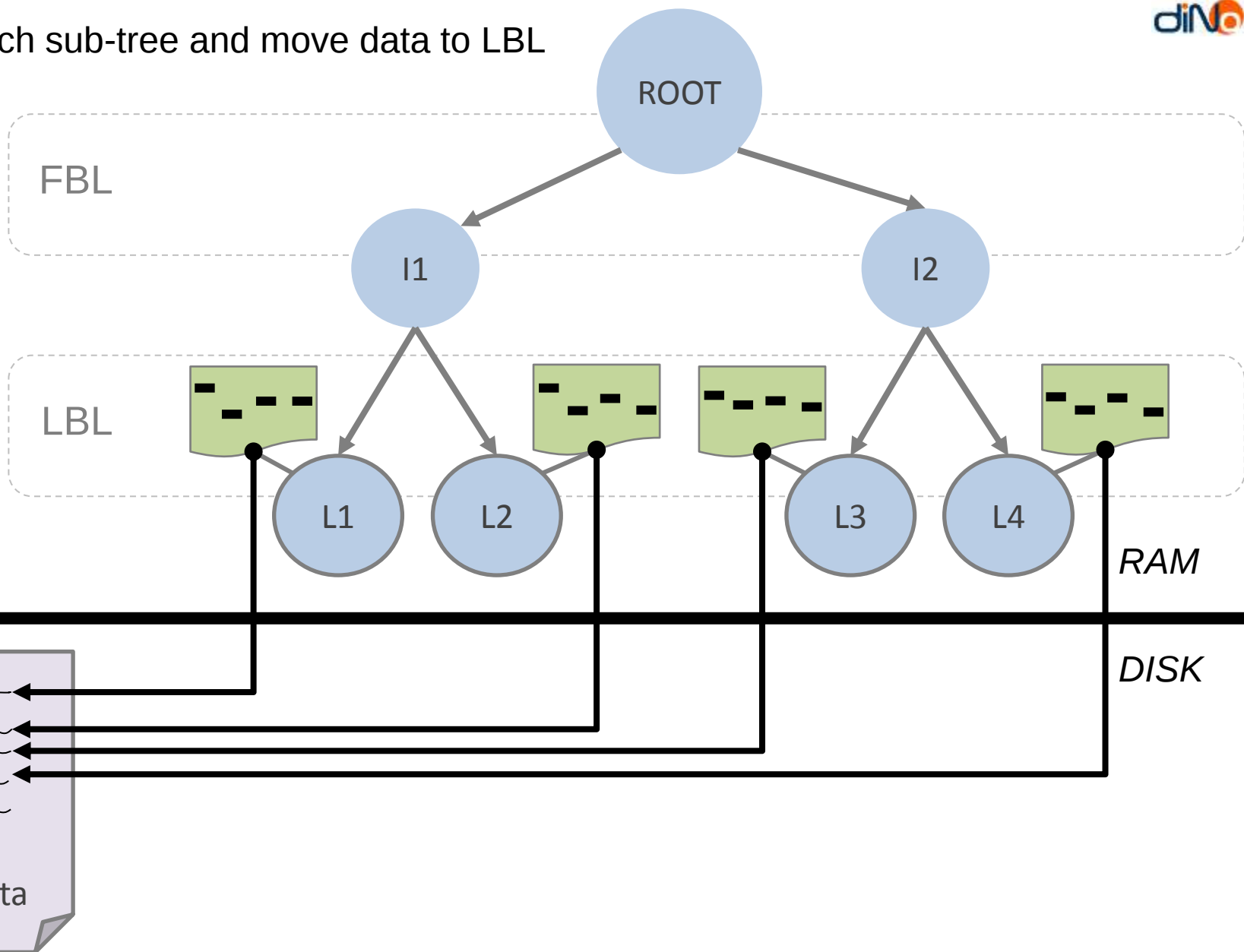
DISK

Discard raw data and keep pointer to raw file

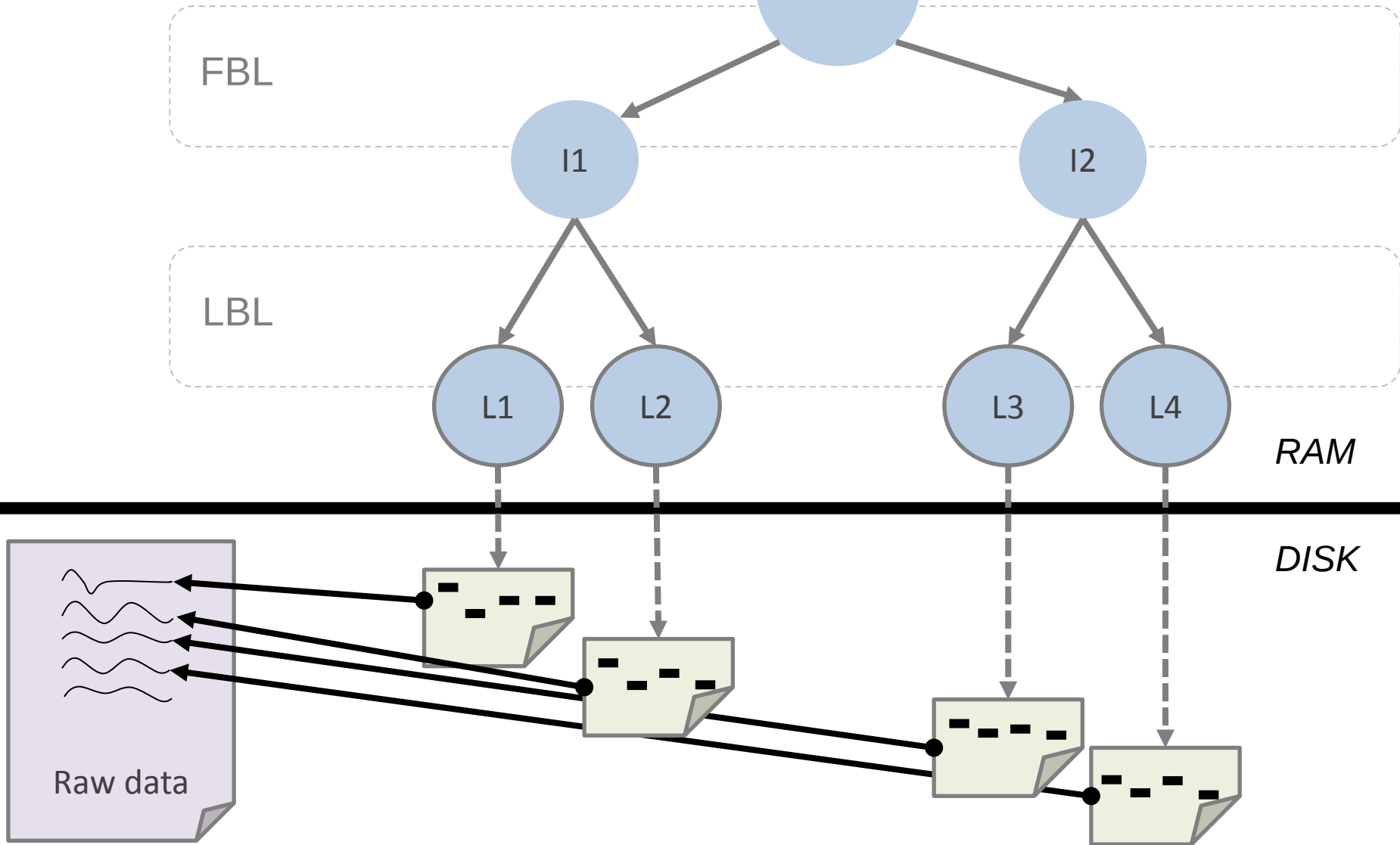


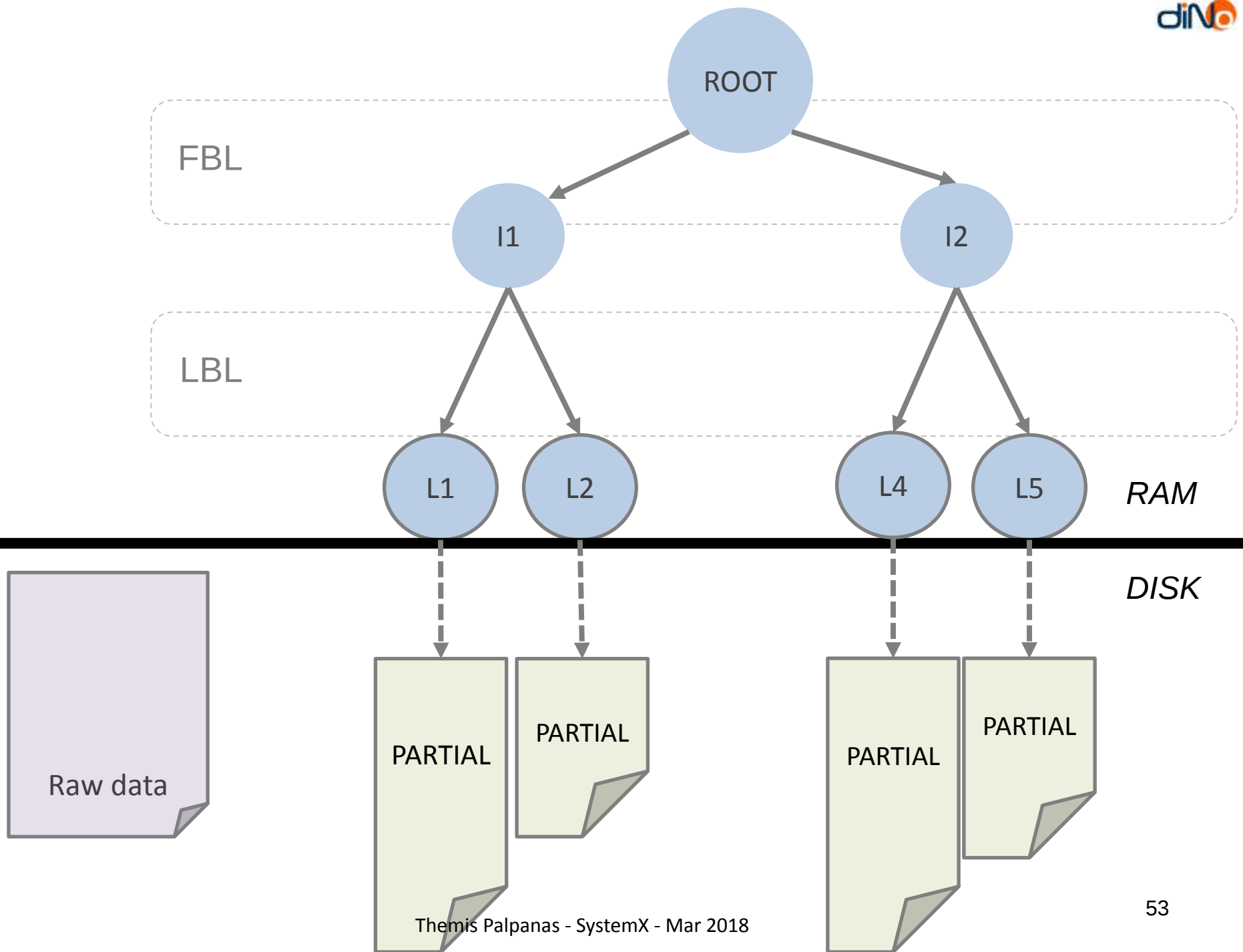


Expand each sub-tree and move data to LBL



We flush the data to the disk to free up memory





Query #1



FBL

ROOT

I1

I2

LBL

L1

L2

L4

L5

RAM

DISK

Raw data

PARTIAL

PARTIAL

PARTIAL

PARTIAL

Query #1



FBL

LBL

ROOT

I1

I2

L1

L2

L4

L5

RAM

DISK

TOO BIG!

PARTIAL

PARTIAL

PARTIAL

PARTIAL

Raw data

Query #1



FBL

ROOT

I1

I2

LBL

L1

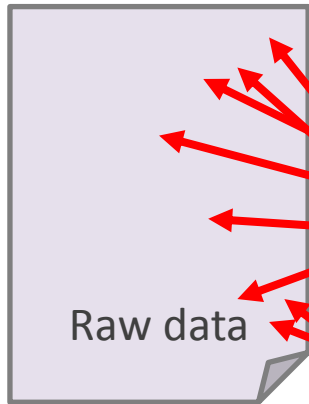
L2

L4

L5

RAM

DISK



Raw data

TOO BIG!

PARTIAL

PARTIAL

PARTIAL

PARTIAL

Query #1



Create a smaller leaf

FBL

Adaptive split

LBL

ROOT

I1

I2

I3

L4

L5

L2

L4

L5

RAM

DISK

Raw data

PARTIAL

PARTIAL

PARTIAL

PARTIAL

PARTIAL

Query #1



Load data in **LBL** and
answer the query

FBL

ROOT

I1

I2

LBL

I3

L4

L5

L2

L4

L5

RAM

DISK

PARTIAL

PARTIAL

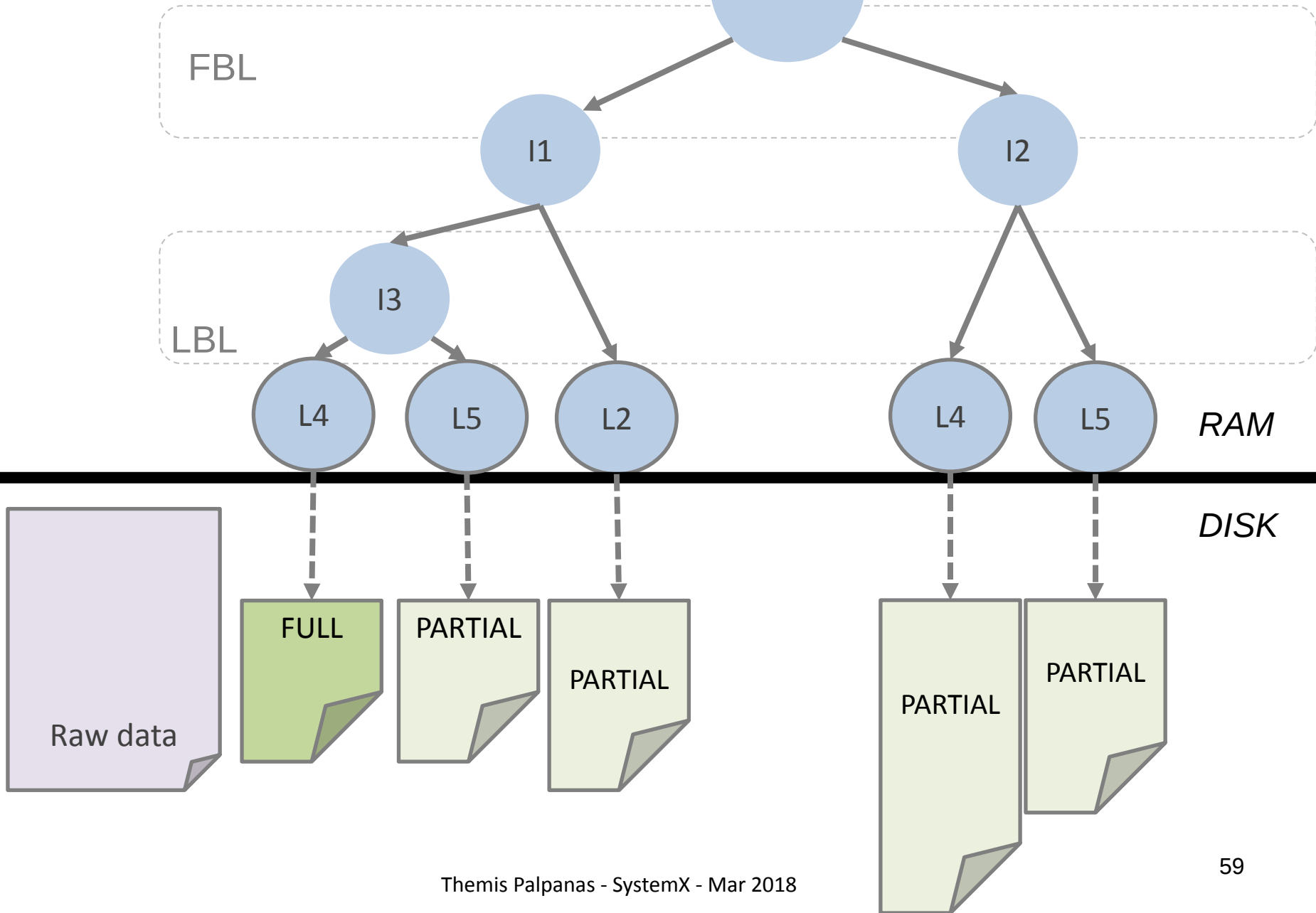
PARTIAL

PARTIAL

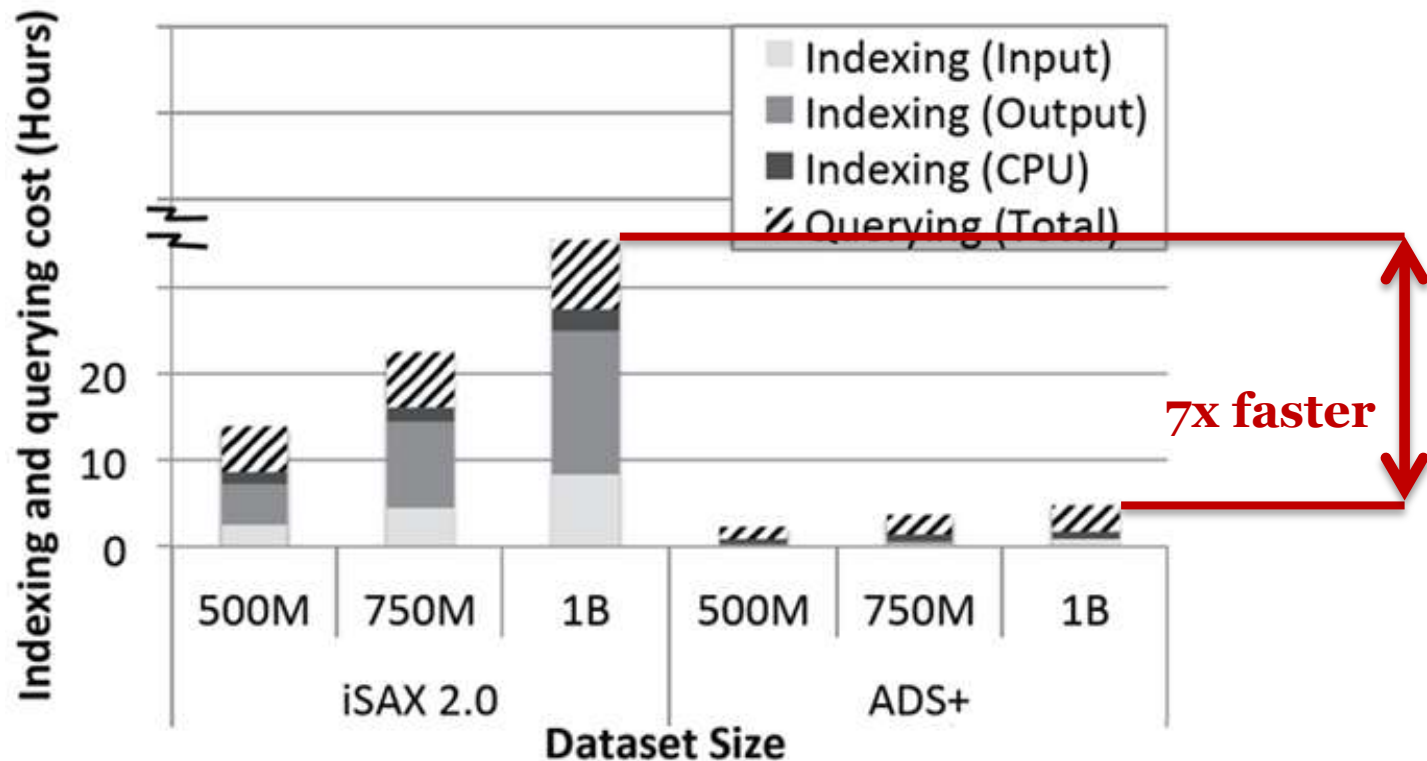
PARTIAL

Raw data

We spill to the disk when we run out of memory



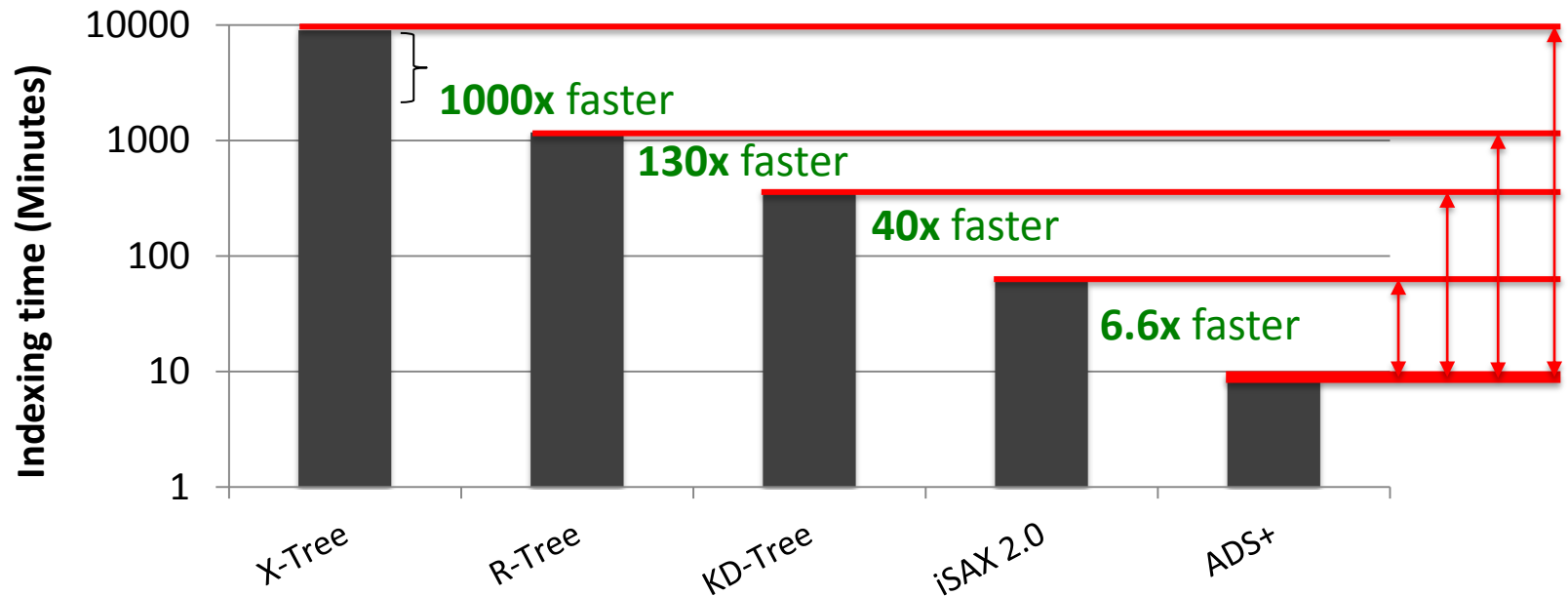
Experimental Evaluation



- iSAX 2.0 needs more than 35 hours to answer 100K approximate queries
- ADS+ answers 100K approximate queries in less than 5 hours

Comparison to *multi-dimensional indices*

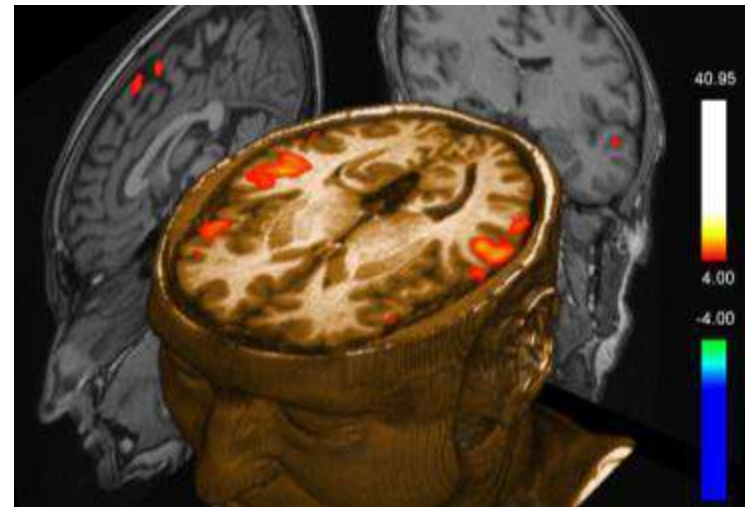
measure data-to-query time
(just index 1 **billion** data-series)



1-3 orders of magnitude faster than multi-dimensional indexing methods

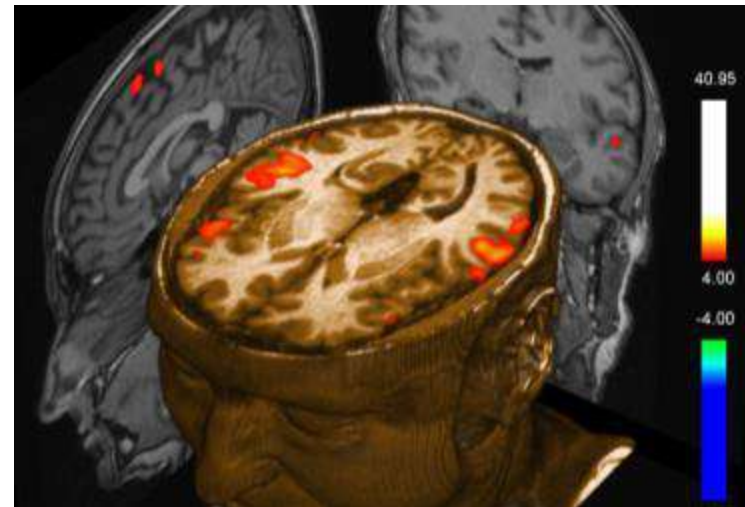
Massive Data Series Collections

- functional Resonance Magnetic Imaging (fMRI) data
 - primary experimental tool of neuroscientists
 - reveal how different parts of brain respond to stimuli
 - single experiment (1 subject, 1 test) produces
 - 60,000 data series of length 3,000: 12 GB



Massive Data Series Collections

- ADHD-200 Global Competition
 - classification task: detect Attention Deficit Hyperactivity Disorder
 - 776 subjects: 9 TB
 - equivalent to: **4.5 billion** non-overlapping data series of size 256
 - equivalent to: **1100 billion** overlapping data series of size 256



Massive Data Series Collections

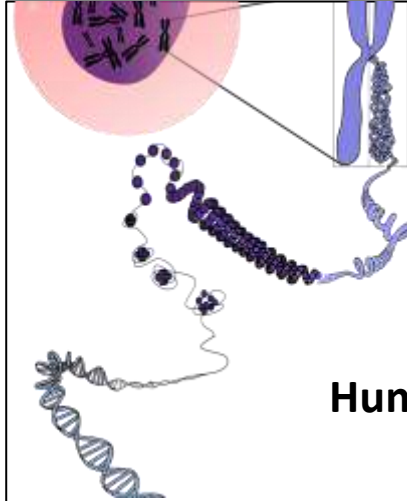


NASA's Solar Observatory

1.5 TB per day

Planned Large Synoptic
Survey Telescope

~30 TB per night



Human Genome project

130 TB



Boeing jet

20 TB per hour

data center and
services monitoring

2B data series

4M points/sec



The Road Ahead

Publications

ICDE'18

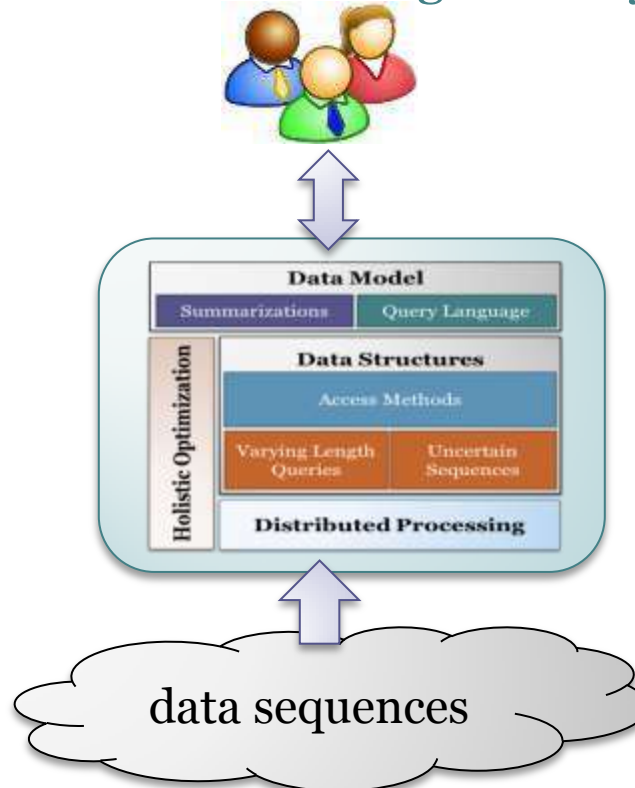
HPCS'17

SIGREC'15

“enable practitioners and non-expert users to easily and efficiently manage and analyze massive data series collections”

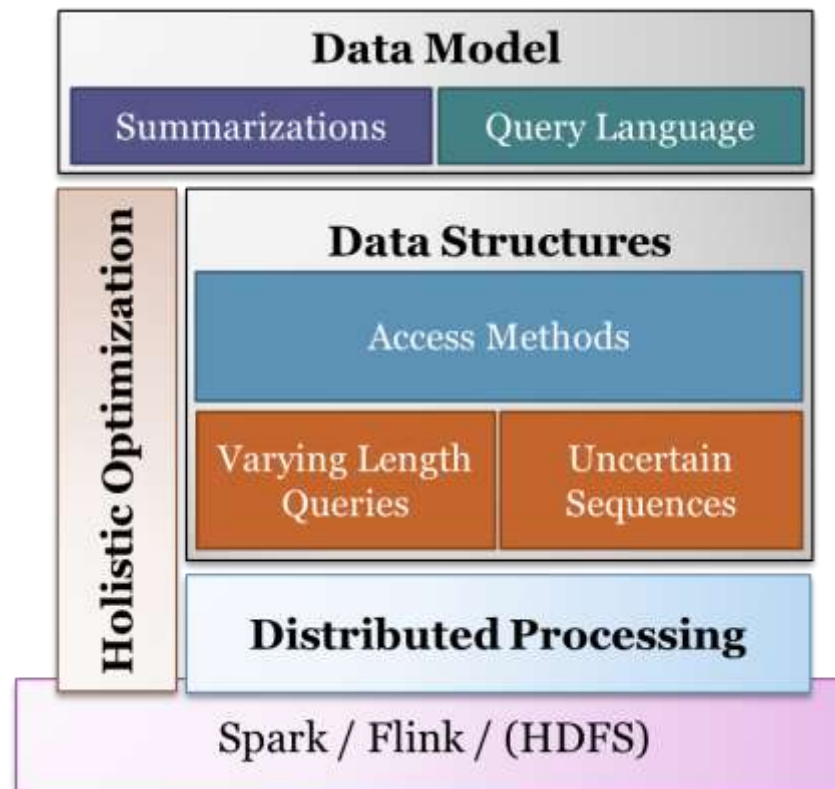
The Road Ahead

- Big Sequence Management System
 - general purpose data series management system



The Road Ahead

- Big Sequence Management System



The Road Ahead

- **Big Sequence Management System**
 - physical and logical data independence
 - rich, declarative query language
 - query optimization
 - inherent scalability
 - support for streaming data series
 - support for uncertain data series

Distribution/Parallelization/Cloud?

- discussion so far assumed serial execution in a single core
 - focus on efficient resource utilization
 - squeeze the most out of a single core
 - produce scalable solutions at lowest possible cost
 - also suitable for analysts with no access to/expertise for clusters

Need for Distribution/Parallelization/Cloud

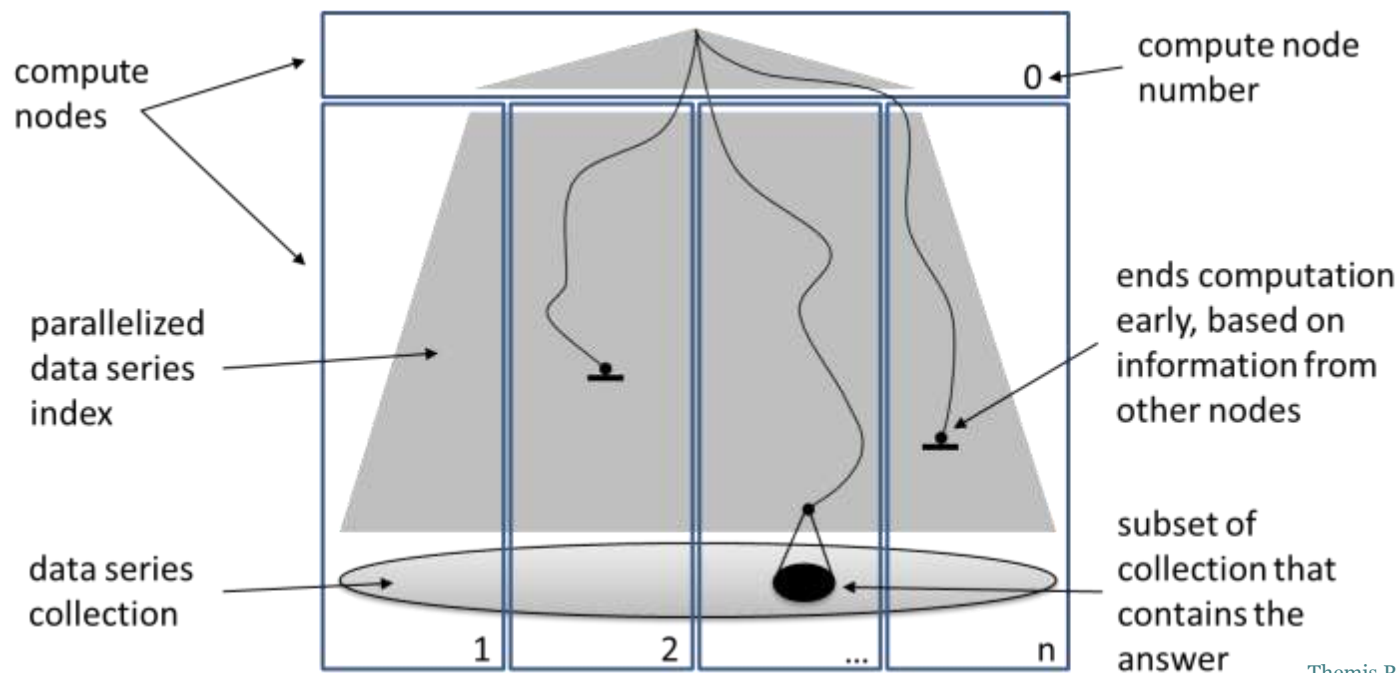
- take advantage of modern hardware!
 - Single Instruction Multiple Data (SIMD)
 - natural for data series operations
 - multi-tier CPU caches
 - design data structures aligned to cache lines
 - Graphics Processing Units (GPUs)
 - propose massively parallel techniques for GPUs
 - new storage solutions: SSDs, NVRAM
 - develop algorithms that take these new characteristics/tradeoffs into account
 - compute clusters
 - distribute operation over many machines

Need for Distribution/Parallelization/Cloud

- further scale-up and scale-out possible!
 - techniques inherently parallelizable
 - across cores, across machines

Publications

ICDM'17



Need for Distribution/Parallelization/Cloud

- further scale-up and scale-out possible!
 - techniques inherently parallelizable
 - across cores, across machines
- more involved solutions required when optimizing for energy
 - reducing execution time is relatively easy
 - minimizing total work (energy) is more challenging

Interactive Analytics?

- data series analytics is **computationally expensive**
 - very high inherent complexity
- may not always be possible to remove delays
 - but could try to hide them!

Need for Interactive Analytics

- interaction with users offers **new opportunities**
 - **progressive** answers
 - produce intermediate results
 - iteratively converge to final, correct solution
 - provide bounds on the errors (of the intermediate results) along the way
 - **imprecise** queries
 - enable user to specify varying accuracy requirements for different parts of the same query
- several exciting **research problems** in intersection of visualization and data management
 - **frontend**: HCI/visualizations for querying/results display
 - **backend**: efficiently supporting these operations

Benchmarking Data Series Indexes?

Previous Studies

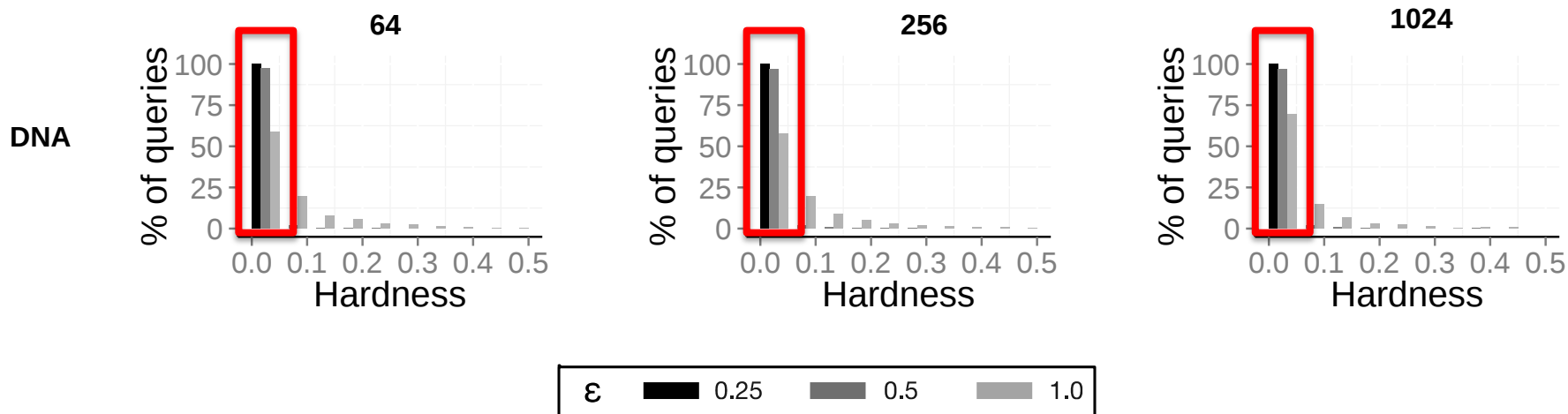
evaluate **performance** of **indexing methods** using **random queries**

- chosen from the data (with/without noise)



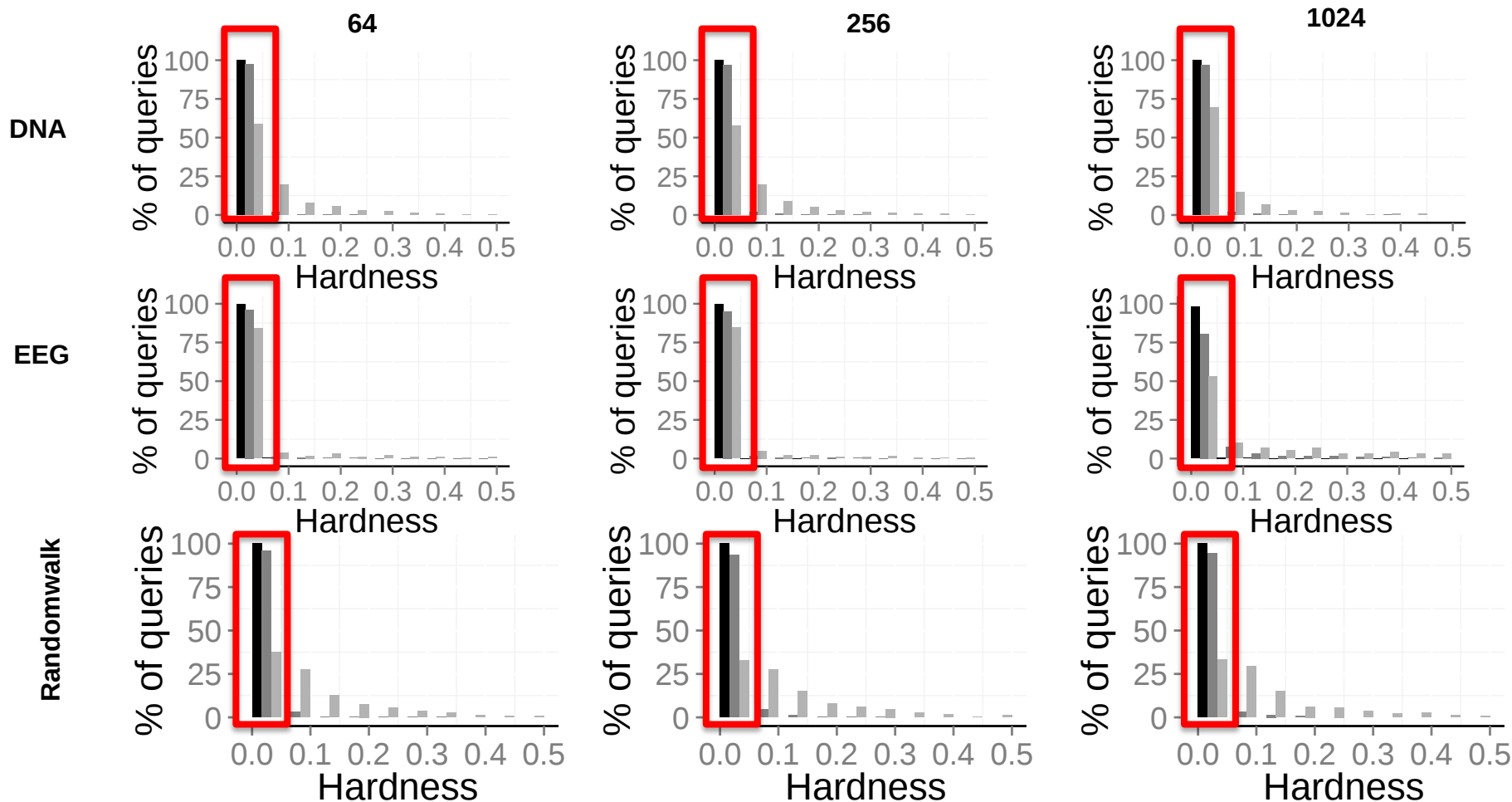
Previous Workloads

Most previous workloads are *skewed* to *easy* queries



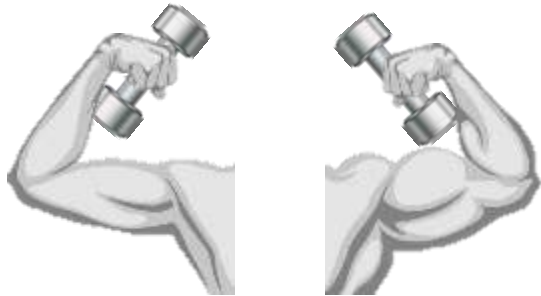
Previous Workloads

Most previous workloads are *skewed* to *easy* queries

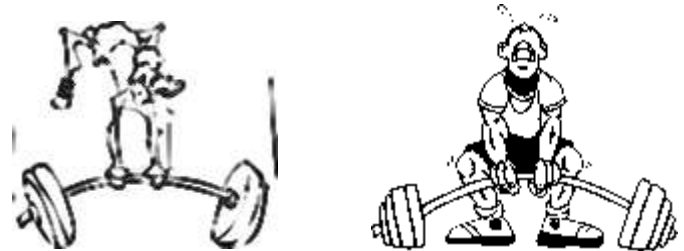


Benchmark Workloads

If all queries are **easy**
all indexes look **good**



If all queries are **hard**
all indexes look **bad**



need **methods** for **generating** queries of **varying hardness**



Publications

KDD '15

Conclusions

- data series is a very **common** data type
 - across several different domains and applications
- complex data series analytics are **challenging**
 - have very high complexity
 - efficiency comes from data series management/indexing techniques
 - all current approaches are ad-hoc
 - waste of time and effort
 - suboptimal solutions
- need for **Sequence Management System**
 - optimize operations based on data/hardware characteristics
 - transparent to user
- several exciting **research opportunities**

collaborations!



Data Series Anomalies Problem

- develop anomaly detection techniques based on sequences (data series), not on individual values
 - individual values can be normal, but their sequence can be abnormal!

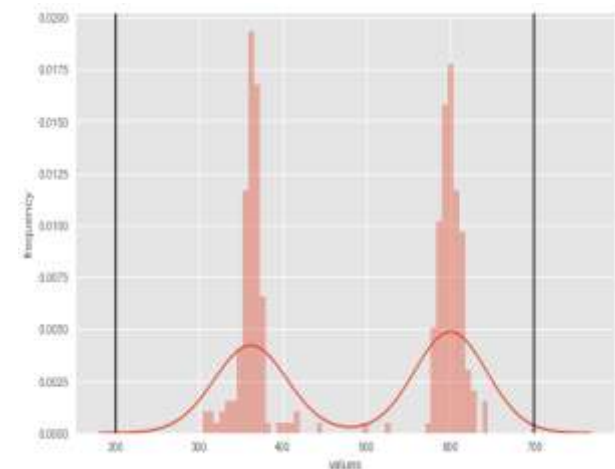
Data Series Anomalies Problem

150 points in a sequence S

- develop anomaly detection techniques based on sequences (data series), not on individual values
 - individual values can be normal, but their sequence can be abnormal!

Minimal critical value

Maximal critical value



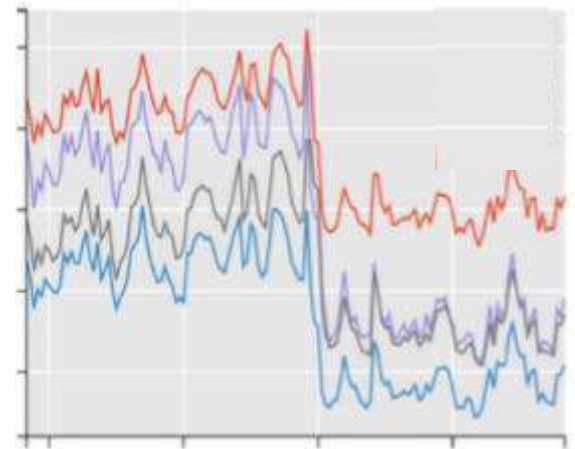
values are not outside critical thresholds

values are normal

Data Series Anomalies Problem

- develop anomaly detection techniques based on sequences (data series), not on individual values
 - individual values can be normal, but their sequence can be abnormal!

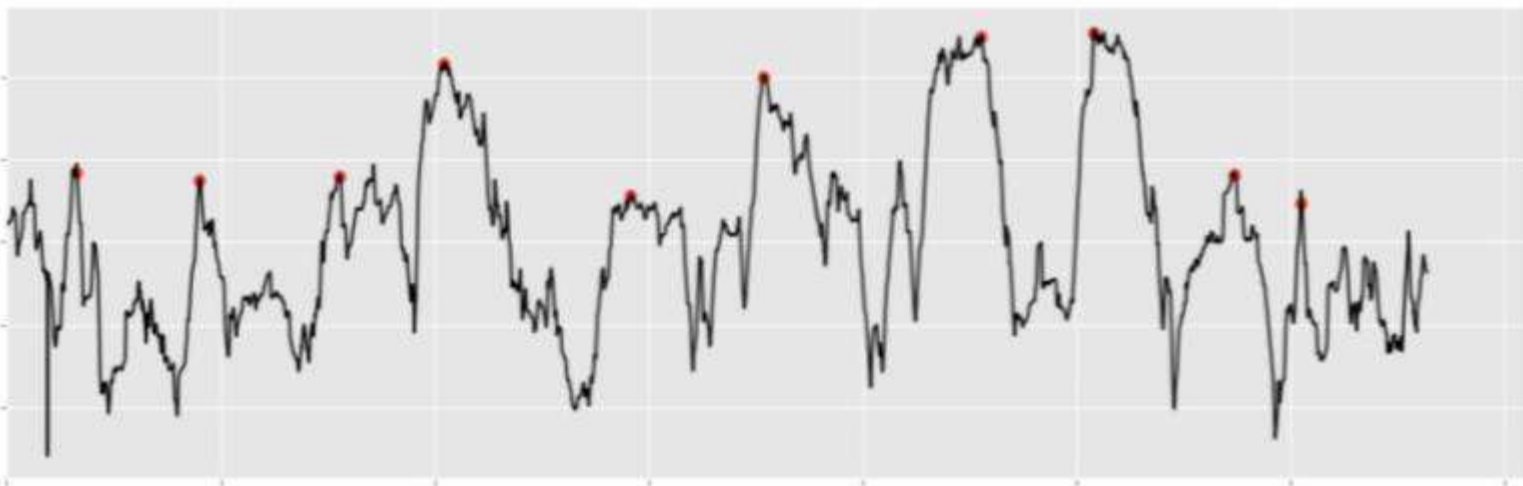
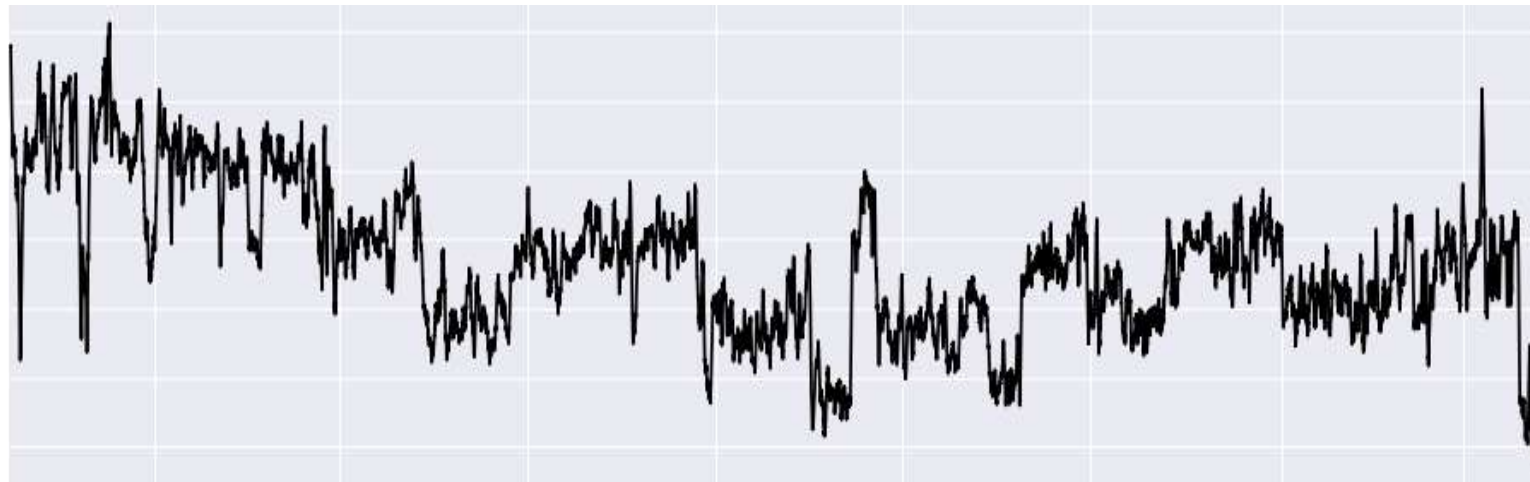
Sequence S



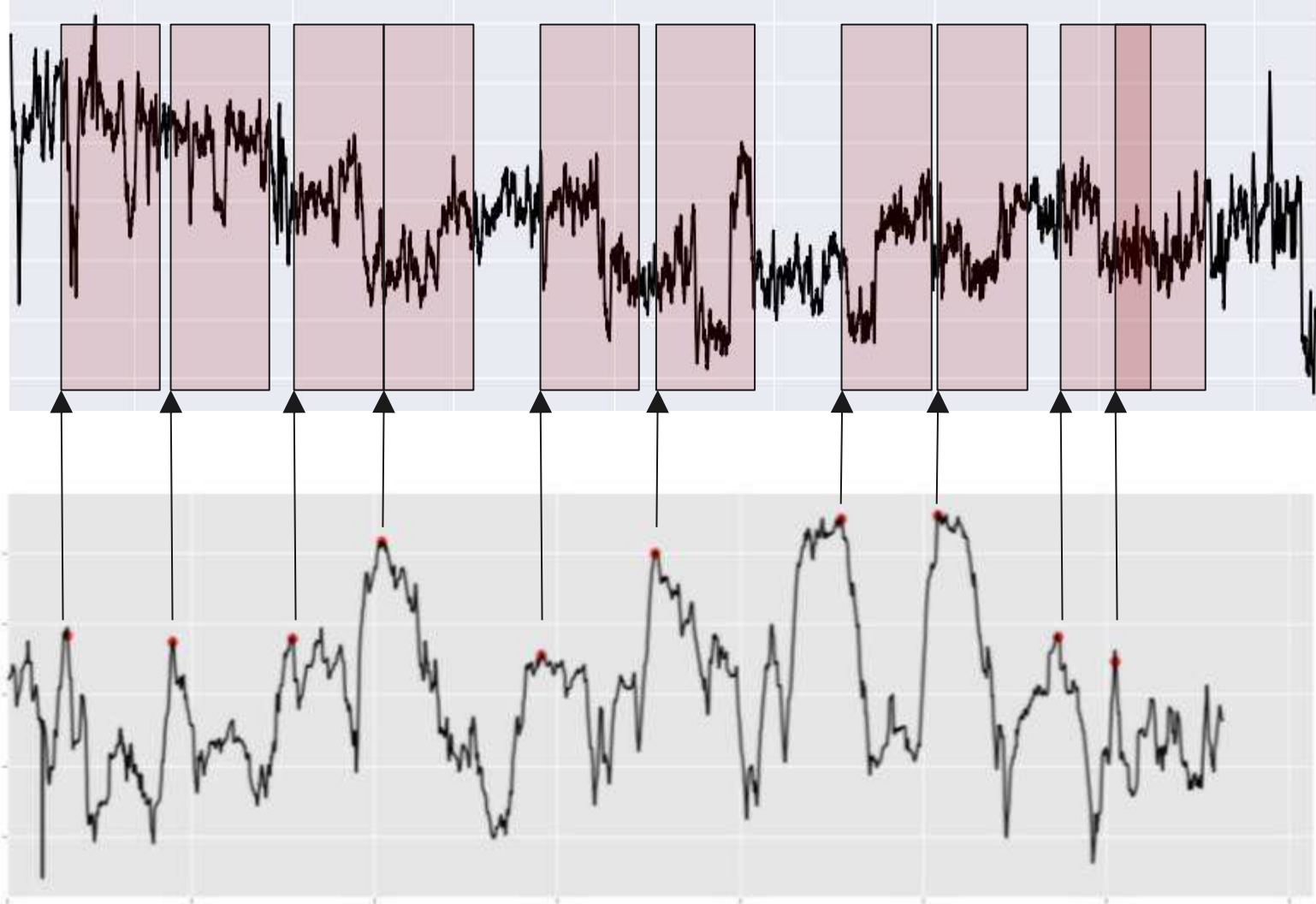
values are not outside critical thresholds

values are normal

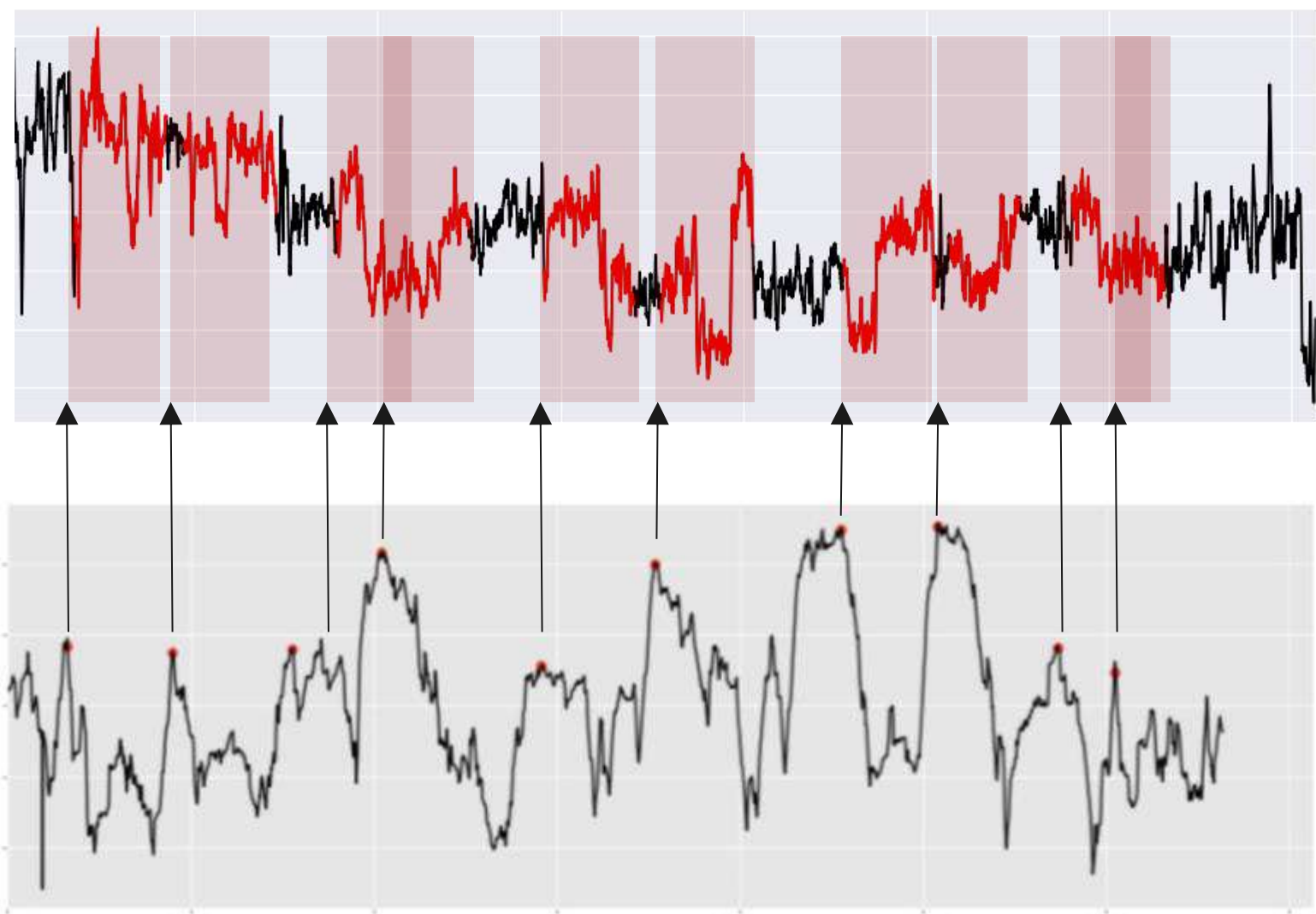
sequences are abnormal



anomalous subsequences



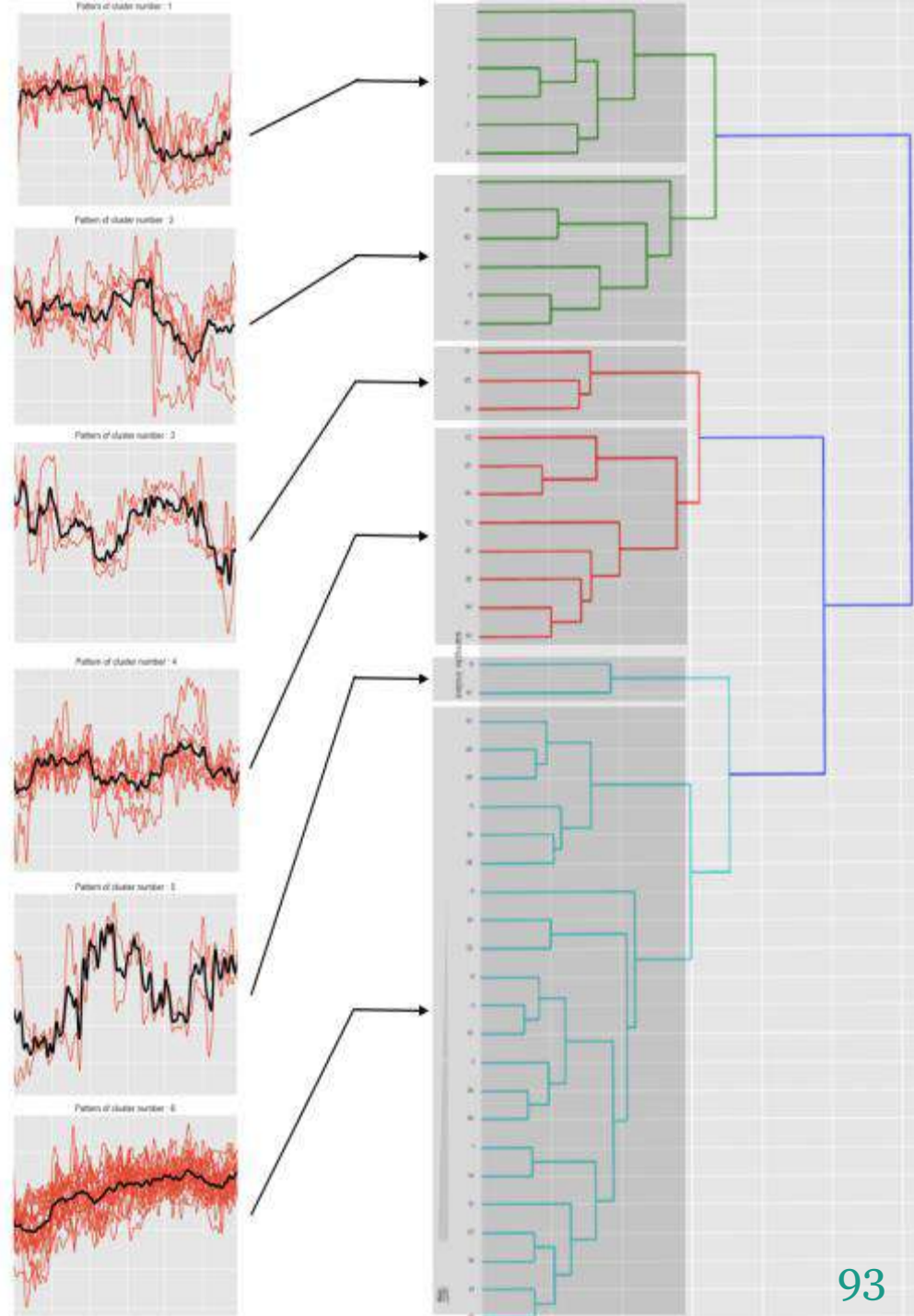
anomalous subsequences



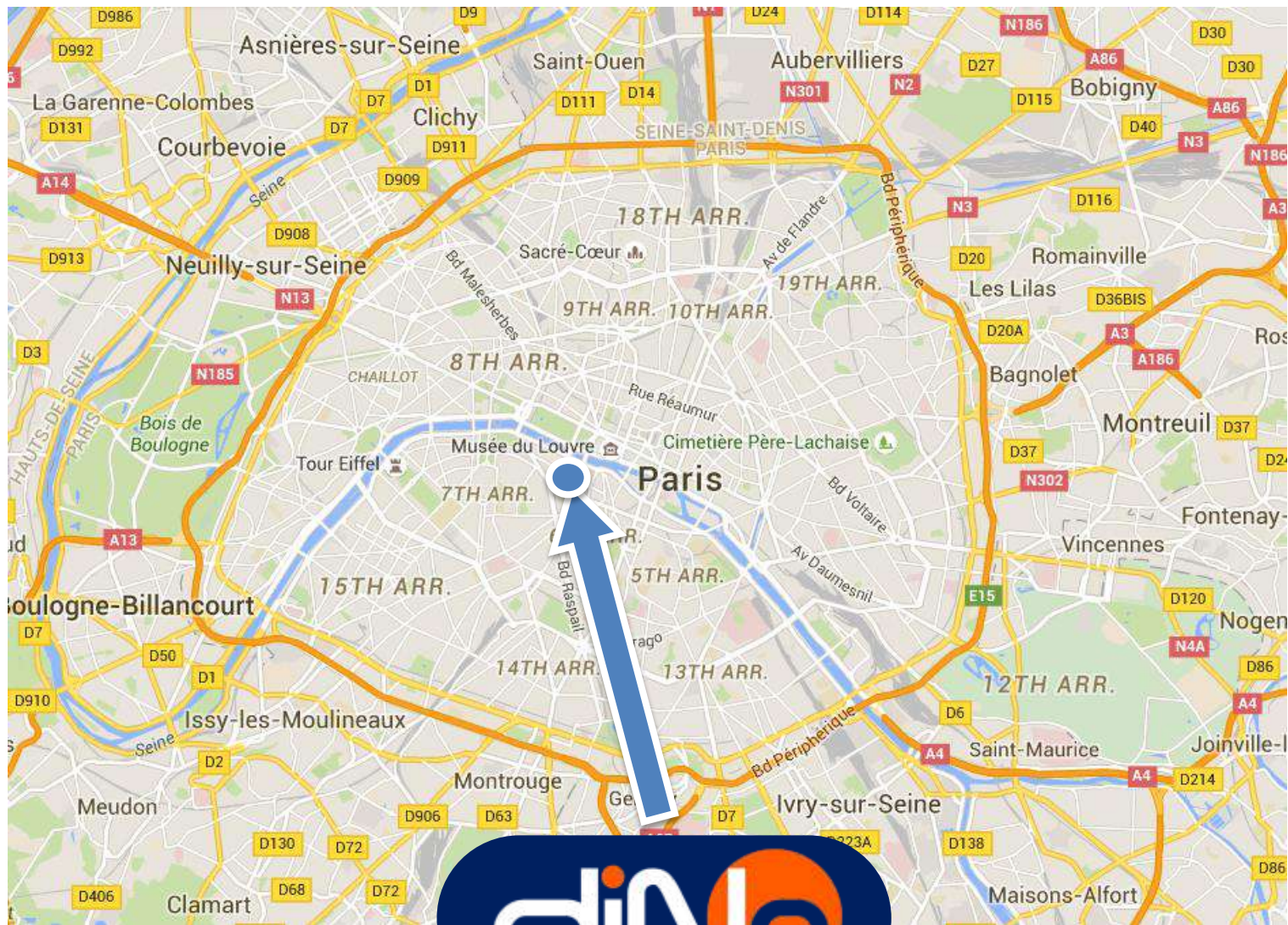
anomalous subsequences

Clustering of anomalies

subsequences grouped
according to their shape



Data-Intensive and Knowledge-Oriented systems



thank you!

google: Themis Palpanas

work supported by:

