

# Balancing Protection and Quality in Big Data Analytics Pipelines \*

Antongiacomo Polimeno,<sup>1</sup> Paolo Mignone,<sup>2</sup> Chiara Braghin,<sup>1</sup> Marco Anisetti,<sup>1</sup>

Michelangelo Ceci<sup>2</sup>, Donato Malerba,<sup>2</sup> Claudio A. Ardagna<sup>1\*</sup>

<sup>1</sup>Dipartimento di informatica, Università degli Studi di Milano, Milan, Italy

<sup>2</sup>Dipartimento di informatica, Università degli Studi di Bari, Bari, Italy

\*To whom correspondence should be addressed;

E-mail: claudio.ardagna@unimi.it.

November 5, 2024

**Abstract:** Existing data engine implementations do not properly manage the conflict between the need of protecting and sharing data, which is hampering the spread of big data applications and limiting their impact. These two requirements have often been studied and defined independently, leading to a conceptual and technological misalignment. This paper presents the architecture and technical implementation of a data engine addressing this conflict by integrat-

---

\*This work was partly supported by *i*) project IMPETUS, funded by the European Union's Horizon 2020 research and innovation programme under grant agreement No 883286, *ii*) project MUSA - Multilayered Urban Sustainability Action - project, funded by the European Union - NextGenerationEU under the National Recovery and Resilience Plan (NRRP) Mission 4 Component 2 Investment Line 1.5: Strengthening of research structures and creation of R&D "innovation ecosystems", set up of "territorial leaders in R&D" (CUP G43C22001370007, Code ECS00000037); *iii*) the project SERICS (PE00000014), funded by the EU - NextGeneration EU under the NRRP MUR program.

ing a new governance solution based on access control within a big data analytics pipeline. Our data engine enriches traditional components for data governance with an access control system that enforces access to data in a big data environment based on data transformations. Data are then used along the pipeline only after sanitization, protecting sensitive attributes before their usage, in an effort to facilitate the balance between protection and quality. The solution was tested in a real-world smart city scenario using the data of the Oslo city transportation system. Specifically, we compared the different predictive models trained with the data views obtained by applying the secure transformations required by different user roles to the same dataset. The results show that the predictive models, built on data manipulated according to access control policies, are still effective.

**Keywords:** Anomaly Detection, Access Control, Big Data, Data Governance, Data Protection

## 1 Introduction

Big data techniques and solutions are pushing towards a data-driven ecosystem where decisions are supported by more accurate analytics and (privacy-aware) data management. The complexity of a data-driven ecosystem is amplified by the heterogeneity and diversity of both infrastructures/tools and data themselves. Data are in fact intrinsically heterogeneous, and collected at high rates and with different formats from different sources in dynamic and collaborative environments.

This data-driven ecosystem comes with two conflicting requirements on data governance. On one side, there is the need to share data in order to extract

value from them, that is, to support fast and accurate approaches for data collection, preparation, analysis. On the other side, there is an increasing need for new solutions protecting data owners from an unregulated data management that could result in a violation of their private sphere. Existing approaches often target a single need independently, aiming to maximize either the quality data analysis by increasing data sharing<sup>1</sup> or the privacy of data owners by favouring data protection<sup>2,3,4</sup>, or either the performance in case of real-time stream processing<sup>5,6</sup>. As a result, there is a lack of proper solutions that balance the two needs, which are often replaced by ad hoc solutions in which each step of the pipeline execution being monitored separately for data governance compliance<sup>7,8,9</sup>.

This paper presents a data engine architecture that supports the execution of big data analytics driven by a new data governance solution based on access control, which ensures a proper data management throughout the pipeline while maximizing the amount of shared data. It extends our work in<sup>10</sup> along four main directions: *i*) the adoption of the access control-based data governance in<sup>10</sup> to enforce data protection requirements at ingestion time, when data are collected and stored, and before data are fed to the analytics pipeline for analysis and processing; *ii*) the definition and implementation of a data engine architecture that supports the enforcement of the access rules; *iii*) the integration of the data engine with modern big data analytics pipelines; *iv*) the evaluation of the solution on an anomaly detection analytics task.

The data governance approach (point *i*)) is built upon the utilization of data annotations and the implementation of secure ad hoc data transformations, conducted during both the data ingestion and processing phases. The secure transformations are selected according to the privileges of the various tasks/users using the data and produce different data views. In this way, our

data engine better balances data protection and usefulness of the data. The result is a platform-independent solution integrated into a novel data engine solution (point *ii*). The proposed data engine architecture introduces granular data control within a big data environment. It has been implemented using Apache tools and components, such as Apache Spark. This implementation guarantees the sufficient scalability for processing the substantial volume of real-time data generated by sensors (point *iii*). We quantitatively evaluated the solution by extending the architecture with a novel state-of-the-art anomaly detection method from sensor network data capable of distributing the workload over multiple computational resources<sup>11</sup> (point *iv*). In particular, the anomaly detector uses different predictive models trained with different data views retrieved according to the privileges of the requesting user and corresponding secure transformations. Then, we evaluated the predictive performance of the different anomaly detection pipelines and the impact of the transformations on the accuracy of the results.

The remainder of this paper is organized as follows. Section 2 describes the motivations and the reference scenario. Section 3 describes our access control-based data governance solution. Section 4 describes the anomaly detection analytics pipelines used during the evaluation phase. Section 5 describes the implementation of our data engine. Section 6 presents the experimental evaluation of our approach. Finally, Section 7 discusses the related work and Section 8 gives our concluding remarks.

## 2 Motivation and Reference Scenario

We present the motivation (Section 2.1) and the reference scenario (Section 2.2) at the basis of the solution in this paper.

## 2.1 Motivation

Big data environments add lots of complexity to data governance and analytics techniques, especially when the need to balance data protection and data sharing arises. Indeed, they are highly dependent on cloud-edge computing, which extensively uses multi-tenancy. Multi-tenancy allows sharing one instance of an infrastructure, platform, or application by multiple tenants to optimize costs. This leads to situations where a service provider offers subscription-based analytic capabilities in the cloud, or the same data engine is accessed by multiple customers. As a consequence, a big data pipeline often entails data and services belonging to different organizations, becoming a serious threat to security and privacy.

Traditional data governance solutions are struggling to keep up with technological evolution for data analytics, and their rigid structure poorly copes with a context where the need for flexibility is one of the fundamental requirements. The emerging conflict between accurate and effective data governance enabling data protection, on one side, and effective data analytics for accurate decision making, on the other side, introduces the need to rethink existing solutions to tune data engine operation to the actual use of the data, as well as to maximize data sharing while providing a suitable level of data protection. New requirements driven by the peculiarities of the big data environment emerged as follows.

**Dynamism and context awareness.** Data are diverse, constantly evolving, and pose new requirements for flexible data governance. Traditional solutions struggle with changing environments and varying user types. User labeling can enhance effectiveness and flexibility. Making such solutions on user labeling rather than static users might contribute to a more effective and flexible approach.

**Ingestion-time enforcement.** Applying data governance at data ingestion permit to reach a better compliance with privacy and data protection laws, avoiding conflicts. By managing and preparing data before storage, transformations can be enforced based on user privileges and processing needs, maintaining compliance, and enabling effective data management.

**Balancing data protection and data quality.** Although data protection and privacy are fundamental requirements when operating in big data environments, they must not clash with the ability to extract as much value as possible from the processed data. An increasing pressure is put on solutions that prioritize the protection of data still preserving a level of data quality and usability.

**Analytics pipeline robustness.** The quality of the data collected is a crucial aspect of any analytics pipeline. In fact, both supervised and unsupervised machine learning methods aim to generate labeled historical data to predict future events. Such patterns and models can produce inaccurate results if data quality is not maintained and information is lost at earlier stages of the pipeline. This is particularly true for unsupervised learning tasks where manually generated labels (i.e., ground truth) are not provided.

**Data analysis explainability.** Data analysis processes must be explainable to each stakeholder. Several applications perform effective analyses without paying sufficient attention to the audience and their ability to understand technical aspects. This is crucial in several applications where the comprehensibility of the responses of the predictive models is fundamental<sup>12,13</sup>.

**Balancing operativity and GDPR compliance.** Nowadays, systems involve

many users with different roles and thus different privileges on the sensitive data to inspect. Therefore, there is the need to maximize the amount of sensitive data and predictive models' output each user can access according to its role. No access should happen only in the worst case scenario.

In this paper, we present an Apache-based big data engine architecture driven by the above requirements, with a novel data governance approach based on an attribute-based access control system<sup>14</sup>. Our data governance approach provides a data governance solution that manages data sharing between different actors of the analytics pipeline with different privileges. The goal of our approach is to balance, while maximizing, quality and protection of shared data and the corresponding analytics. We integrated the data engine with a smart data analytics solution for anomaly detection and experimentally evaluated our approach in a smart city reference scenario.

## 2.2 Reference Scenario

Our reference scenario considers a smart city scenario, where public transportation data in the city of Oslo (i.e., buses and trams) are used for anomaly detection in traffic patterns. Data are collected in real-time and aggregated on their position every 5 minutes. Aggregated data are ingested in batches in the data storage. Each batch contains information about the GPS location (latitude and longitude) of the vehicles, the status of the vehicles (malfunction, traffic congestion) and details about the path (origin, destination, stop point, delay).

The transportation data are used by a traffic monitoring application (i.e., anomaly detector) that comes in three configurations, one for each category of users: policeman, air quality control officer, and citizen. Each user has different access rights applying different data transformations, since not all users can access the whole dataset. For instance, local policy department can use the

data to detect specific locations in which intervention is needed, whereas the local air quality department can use the data to detect areas with spikes in traffic, possibly helping citizens avoid congested areas.

Each configuration of the anomaly detector consists of a different machine learning pipeline built on a different machine learning model. All models are derived from the same dataset at different levels of granularity due to the applied data transformations. A model trained on “normal”, real-time traffic data of 5-minute batches represents the usual traffic behavior and is used to spot any divergences from the expected behavior. Indeed, once collected and ingested, data are prepared for the analysis with traditional data preparation (e.g., cleaning). Then, data are transformed according to a set of applicable access control policy rules that evaluate the intended user and the requested data. Finally, data are used by the anomaly detector.

To this aim, a cloud-based multi-tenant data engine is implemented to support our anomaly detector, and in general third-party smart city applications, to work on the same datasets. The data engine is extended with data governance approach built on an attribute-based access control system. Considering data analytics, our anomaly detector assumes data to be shared for data analysis. This is architecturally and methodologically different from the approach used in federated learning<sup>15</sup>, where data are generated and used locally at multiple nodes to create models, whose parameters are shared while retaining raw data private.

In order to guarantee explainability, the proposed trained model can distinguish between normal and anomalous cases by labeling the current data gathered by sensors. The algorithm also provides the “motivations” of an anomaly, explaining the raised alerts to the security operators. The motivations are implemented as a ranking of the variables under analysis in descending order w.r.t.



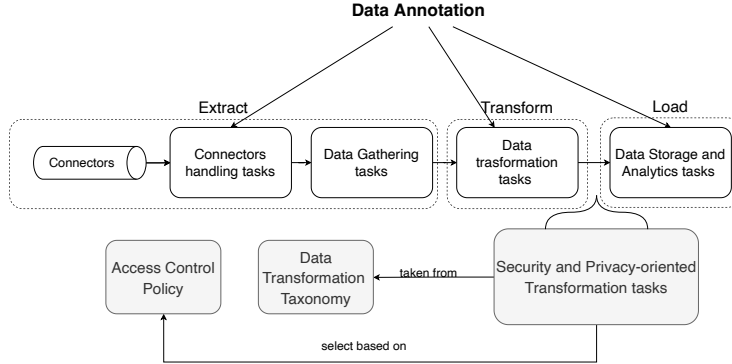


Figure 1: Annotations transformations and access control policies applied to our ETL ingestion pipeline

their importance (e.g., for a bus stuck in traffic jam, the average delay of vehicles within the involved zone will appear as top-ranked, while the level of smog in the air will appear as the second-ranked, and so on).

This approach is motivated by the reference scenario where it is requested that the predictive models must explain the urban problems since police operators can promptly understand their type and develop appropriate measures and that must be suitable for the type of problem encountered.

### 3 Data Governance

Our policy-aware data ingestion procedure<sup>10</sup> for data governance implements an advanced Extract, Transform, Load (ETL) approach. It extends traditional ETL with data transformations triggered by access control policies and is used to sanitize ingested data. These policies are based on the identity of the user/service using data and the annotations on the ingested data.

### 3.1 Access Control-Based ETL

Figure 1 shows how a traditional ETL is extended for better data governance. The white boxes in Figure 1 refer to a traditional ETL ingestion procedure in the form of a big data pipeline that includes: i) *connectors* establishing the connection with the data sources extracting raw data with producer speed; ii) *connectors handling tasks* dealing with the different connectors’ peculiarities, e.g., raw data format/size and the velocity of data producers; iii) *data gathering tasks* collecting data; iv) *data transformation tasks* transforming data formats; v) *data storing/forwarding tasks* either saving or forwarding data to the analytics procedures.

Data annotation on top of Figure 1 annotates data at different steps of the pipeline as follows: i) connection handling tasks, to annotate raw data at gathering time; ii) data transformation tasks, to annotate specific data fields before their transformation; iii) data storing and analytics tasks, to provide annotations on data resulting from the transformation tasks. The grey boxes in Figure 1 extend the traditional ETL ingestion pipeline (white boxes) to address the security and privacy challenges of multi-tenancy scenarios with data transformation processes built on access control policies and data annotations. Security and privacy-aware transformations are used, ranging from pruning and reshaping to encrypting/decrypting or anonymizing the full resource or part of it. These transformations are selected based on a taxonomy of data transformations and triggered by access control policies, leading to a more dynamic and flexible ETL ingestion procedure.<sup>1</sup> This approach makes it possible to decouple the data transformations typical of a traditional ingestion procedure from the policy-driven data transformations we propose in this paper, which are linked to a specific processing task and the requesting user/service.

---

<sup>1</sup>We note that the transformations triggered by access control policies can be implemented as big data pipelines themselves.

### 3.2 Access Control Policies

The access control language proposed in this paper is a refinement of our previous work in<sup>10</sup>. It extends an attribute-based access control model that offers flexible fine-grained authorization capabilities and exploits XACML<sup>16</sup> obligations to introduce preemptive data transformations. Indeed, in a collaborative big data scenario, preemptive data transformations are more suitable than denying access to data. For this reason, when a subject makes an access request to a resource, our access control filters the set of results according to the subject's privileges, obtaining data protection by removing or obfuscating sensitive attributes rather than denying access to full data.

We tailor the common key elements of an attribute-based access control model to deal with the peculiarities of a big data scenario. The access control policy expresses access conditions based on key elements and their attributes as follows.

**Definition 3.1** (Policy Condition). A *Policy Condition* is a Boolean expression of the form  $(attr\_name \text{ op } attr\_value)$ , with  $op \in \{<, >, =, \neq, \leq, \geq\}$ ,  $attr\_name$  an attribute label, and  $attr\_value$  the corresponding attribute value.

A policy is then defined as follows.

**Definition 3.2** (Policy). A *policy*  $P$  is 5-uple  $\langle subj, obj, action, env, data-trans \rangle$ , where:

- Subject  $subj$  defines a user or the service provider of a job issuing access requests to perform operations on objects. It is of the form  $\langle id, PC \rangle$ , where  $id$  defines a class of users (e.g., policeman), and  $PC$  is a set of *Policy Conditions* on the subject, as defined in Definition 3.1. For instance,  $\langle user, \{age \geq 18\} \rangle$  refers to a person of legal age.
- Object  $obj$  defines any data whose access is governed by the policy. It is of

the form  $\langle type, PC \rangle$ , where: *type* defines the type of object, such as a file (e.g., a video, text file, image, etc.), a SQL or noSQL database, a table, a column, a row, or a cell of a table, and *PC* is a set of *Policy Conditions* defined on the object's attributes. For instance,  $\langle file, \{(creation\_date < 2020/12/05)\} \rangle$  refers to a file created before 2020/12/05.

- Action *action* defines any operations that can be performed within a big data environment, from traditional atomic operations on databases (e.g, CRUD operations varying depending on the data model) to coarser operations, such as an Apache Spark Direct Acyclic Graph (DAG), an Hadoop MapReduce, an analytics function call, or an analytics pipeline.
- Environment *env* defines a set of conditions on contextual attributes, such as time of the day, location, IP address, risk level, weather condition, holiday/workday, emergency. It is a set *PC* of *Policy Conditions* as defined in Definition 3.1. For instance,  $\{(weather\_conditions = red\_alert)\}$  refers to a red alert broadcasted for weather conditions.
- Data Transformation *datatrans* defines a set of security and privacy-aware transformations on *obj*, focusing on compliance to regulations and standards, in addition to simple format conversions.

More specifically, the data transformations in the policy include, but are not limited to, suppression or generalization of data, masking and encryption, distortion, and swapping. They are grouped in a taxonomy of functions used to sanitize the ingested data according to the security property they aim to guarantee. We consider the CIA triad defined by NIST<sup>17</sup> as the reference for our work as follows.

- *Confidentiality* describes the process of hiding information that may lead to personal identification or proprietary information, keeping users anony-

mous. The conventional security mechanisms used to protect data confidentiality are encryption or hashing<sup>18,19,20</sup>.

- *Integrity* ensures that data are non-repudiable, authentic and protected from modifications and deletions. Some of conventional security mechanisms used to ensure data integrity are encryption and signing<sup>18 19 20</sup>.
- *Availability* ensures that authorized parties should have constant and easy access to information. This includes correctly maintaining hardware, technological infrastructure, and systems that store and show data. Some of conventional security mechanisms used to ensure availability are replication, distribution, high availability and disaster recovery.

Data transformations can be further classified on the type of information that is sanitized. *Temporal transformations* delay access to data. Data produced at a time  $T$  will be accessible at a time  $T + \delta$ , with  $\delta$  depending on the policy applied and the desired level of visibility. For example, as shown in Figure 2, a Policeman accesses data when produced (green segments), while an air quality control officer retrieves them with 1-hour delay (red segments).

*Spatial transformations* alter the granularity of geolocation information, making it more or less granular based on the applied policy and desired level of visibility. For example, in the smart city context, if an unexpected event occurs (e.g., fire, accident), law enforcement could have a more spatially accurate view, while citizens could be shown only the macro area where the event occurred.

### 3.3 Access Control Policy Enforcement

When an access request is submitted, the set of applicable policies is first selected and then enforced on the target object. First of all, the access request is evaluated against the *subj*, *obj*, *action*, and *env* in  $P$ . If the conditions *cond*



Figure 2: Example of Temporal Transformation

in *subj*, *obj*, and *env* are positively evaluated according to the access request, and the actions in the policy matches the action in the request, the policy is enforced. Policy enforcement consists of the execution of *datatrans* in *P* against the object target of the access request.

In the following, we present two examples of policies in the reference case study we presented in Section 2.2.

- **Policy 1 (No Transformation).** This policy guarantees unrestricted access to data (no transformation). It is triggered when a *local police department* officer (subject), in an *ordinary* context (*env*), requests to *read* (action) data tagged as *PII-id* (object). It can be defined as follows:

```

subject = (user,{(department="LocalPolice")})
object = (file,{(tag="PII-id")})
env = (traffic conditions="Ordinary")
action = read
datatrans = ∅

```

- **Policy 2 (Temporal Transformation).** It performs a *temporal transformation* granting access only to data older than 24 hours. The policy is triggered when an *Air Quality Department* officer, in an *ordinary* context, requests to *read* data tagged with *PII-id*.

```

subject = (user,{(department="AirQuality")})
object = (file,{(tag="PII-id")})

```

```

env = (traffic conditions="Ordinary")

action = read

datatrans = {TimeShift_1D}

```

- **Policy 3 (Spatial Transformation).** It performs a *spatial transformation* aggregating data. The policy is triggered when a *Citizen*, in an *emergency* context, requests to *read* data tagged with *PII-id*.

```

subject = (user,{(department="Citizen")})

env = (traffic conditions="Emergency")

action = read

datatrans = {LowGranularityAggregation_1D}

```

## 4 Data Analytics for Anomaly Detection

Anomaly detection is one of the most important tasks within a smart city reference scenario, where several sensors produce data that could be continuously collected and monitored. In this paper, we adopt a machine learning method based on deep neural networks to perform anomaly detection in a distributed fashion. Our method explains the detected anomalies and works in an unsupervised setting.

### 4.1 Anomaly Detection

Anomaly detection is a machine learning task that aims to learn a predictive function by exploiting the historical data representing normal cases. The ability to distinguish normal and anomalous cases in data is crucial to create good predictive models. In this paper, we build a resilient model based on Spark-GHSOM<sup>11</sup> that is robust to noisy data representing (usually) few anomalies within the training set. Furthermore, Spark-GHSOM<sup>11</sup> is a state-of-the-art distributed method for learning predictive models. It is mainly designed for

multi-target regression and forecasting problems. We integrate Spark-GHSOM into the full architecture to perform explainable anomaly detection, where explainable predictions are provided to the final users. The process for training the predictive models of the Spark-GHSOM follows the classical process of the GHSOM (Growing Hierarchical Self Organizing Maps) training<sup>21,22</sup> and is enhanced to work properly also for mixed heterogeneous attributes<sup>23</sup>. It also exploits the Apache Spark framework via the Map-Reduce paradigm for efficient and distributed learning.

The first step in the GHSOM algorithm computes the inherent dissimilarity in the input data with different types of attributes. For numerical attributes, the method exploits the *mean quantization error*, while for categorical attributes the *unlikability* represents a good measure to estimate the difference among values<sup>24</sup>. The method uses such dissimilarity functions to construct SOMs, where a single SOM represents a neural network composed of a grid of neurons, usually at the starting point of a  $2 \times 2$  map of neurons, which can be enlarged, if needed, at the training phase.

During the training phase, a procedure modifies each neuron that is not sufficiently descriptive of its surrounding training examples in order to enhance the quality of the current model<sup>11</sup>. Therefore, such a neuron is substituted by another SOM in a hierarchical structure of SOMs (see Figure 3) and the training proceeds recursively.

After the training stage, the set of neurons within the GHSOM are used to identify, for each new unseen instance, its closest neuron within the hierarchy of SOMs. Such a neuron represents the model part that best fits and represents the input instance. If such an instance is relatively close to the identified neuron, then the input instance is a normal case; otherwise, the input instance is considered an anomaly.



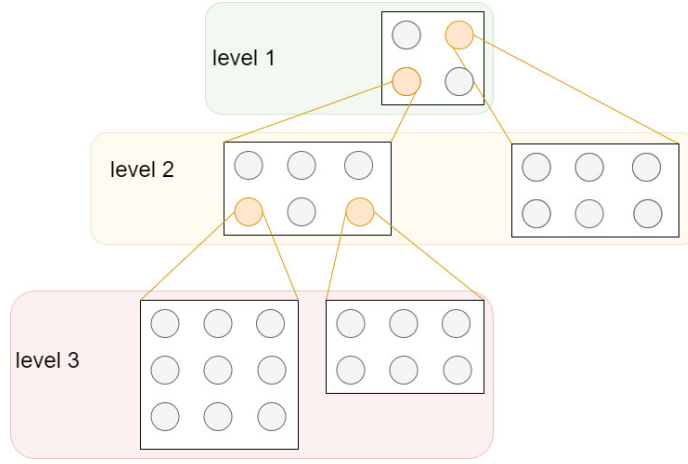


Figure 3: A Growing Hierarchical Self Organizing Map.

Formally, let  $x$  be the new instance to be considered and  $n$  one of the possible neurons within the GHSOM, and let  $n(x) = \arg \min_n \text{dist}(x, n)$  the closest neuron to  $x$ . The instance is considered an anomaly if  $\text{dist}(x, n(x)) > (\mu + t\sigma)$ , where  $\mu$  represents the average distance of the training instances and the neurons of the model,  $\sigma$  the standard deviation of such distances, and  $t$  the user-defined threshold factor parameter.

## 4.2 Explainability

Neural networks such as SOMs are good predictive models but with low explainability, often characterized by the impossibility of understanding the motivations of their output. In several real applications, the need of explainability is high. For instance, in a smart city scenario, a security operator should be able to understand why an anomaly is raised by a system in order to act accordingly.

Our extended version of SparkGHSOM is capable of explaining the detected anomalies. Indeed, a more informative approach is considered combining the Boolean response (normal/anomaly) with a ranking of the features describing the detected anomalous phenomenon. Therefore, it is possible to produce a

feature ranking according to their importance, indicating their contribution to the detection of the potential anomaly. More specifically, a feature ranking produces a ranking, according to the feature importance measure, of the whole set of features describing the whole set of historical instances considered for training the predictive model, while the feature importance is quantified through a number between 0 and 1 indicating how anomalous the value expressed by the feature is w.r.t. the data collection (the total sum up to 1).

The importance score is determined starting from the Euclidean distance between the current test instance under analysis and its closest neuron within the GHSOM model. More formally, the method computes a ranking w.r.t. the contribution provided by the single feature in the distance between  $x$  and  $n(x)$ . The ranking function for the instance  $x$ , is computed as  $r(x) = \frac{(x[i]-n(x)[i])^2}{\sum_j (x[j]-n(x)[j])^2}$ , where  $i, j$  represent feature indexes. This approach indicates the prominent feature(s) to describe the raised alert explaining its reason. For example, in urban scenarios, for a specific area and time, an event could be described through the level of traffic, the concentration of pedestrians, the level of air pollution, the local temperature, the road visibility level, and so forth. Supposing that a detected anomaly comes with an indication that the traffic level and air pollution level are unusually high in and around a specific location, the operator can gain insights into the nature of the anomaly which in this example could be associated with a fire in a building facing the street.

## 5 Data Engine Implementation

Figure 4 shows the architectural view of our data engine enforcing access control before data is used. The data engine consists of an ecosystem of data management tools and components that support the management of the entire data life-cycle, from ingestion to analytics. It defines the whole data management

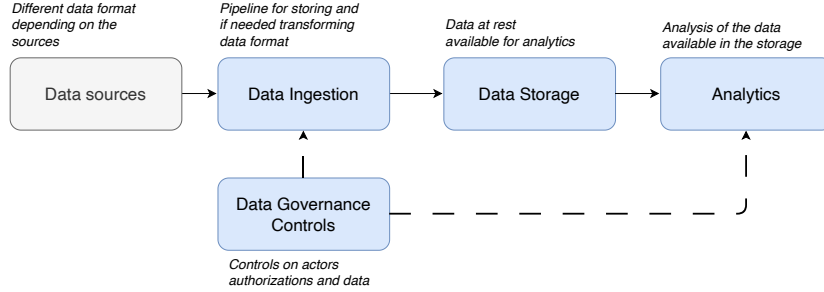


Figure 4: Architectural view of the data engine

process, that is, the process of collecting, storing, using, and maintaining data in a secure, efficient, and cost-effective way. It includes all the practices devoted to the protection of data, as well as all the activities needed to prepare properly data for an analytics process.

The data engine acts as a middleware between *data sources*, i.e., physical data collection points, and *analytics*, i.e., the pipeline that implements a specific analysis on the data available in the data storage to extract value, forming a complete big data engine.

The solid black arrows in Figure 4 show the standard flow of data from data sources to analytics. First, the data are collected from a variety of data sources, each of which has a possible different format. The data are then forwarded to the ingestion pipeline that specifies the activities needed to transform and reconcile the different data formats. Finally, the ingested data are stored in a data store for further analytics activities. The dashed arrows highlight that the whole data management process is subject to data governance controls, where ad-hoc controls are implemented to monitor actors' activities, data exchanges and operations. In particular, access control enforcement is executed both during data ingestion and before any data processing task to manage access to data

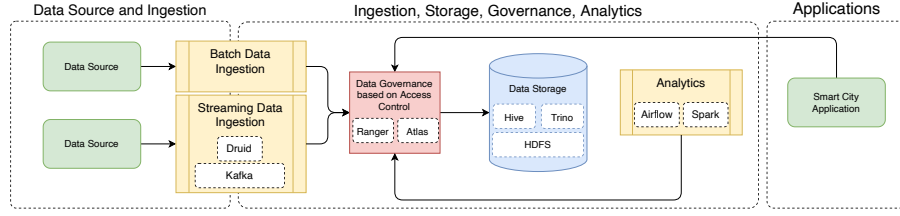


Figure 5: Apache-Based Big data engine.

prior to its storage and manipulation.

We implemented the data engine in Figure 4 using services and components of the Apache-based ecosystem, integrating the data governance approach in Section 3 and the data analysis approach in Section 4 within a complete Apache-based Big Data engine. Figure 5 highlights all the building blocks of the Big Data Engine, their interactions, and the technologies on which they are based, as follows.

**Data ingestion** The data ingestion procedure involves storing data in the data storage, which may require format transformations. It is implemented as a pipeline with tasks for data collection, transformation (if necessary), and saving to the storage. Data ingestion supports batch and streaming modes using our data governance approach based on access control.

In batch ingestion, data is collected either through the source API or manual upload, with options for active (event-based) or passive (time-based) data collection. *Apache NiFi*, *Hive*, *Trino*, and *HDFS* are used for batch ingestion.

For stream ingestion, *Apache Kafka* and *Druid* are used. *Druid* enables real-time transformations on data streams and works together with *Kafka* for ingestion-based access control enforcement on stream data.

**Data storage** The data storage component is responsible for storing data while

ensuring confidentiality, integrity, and availability. It provides data to analytics and includes features for data availability, fault tolerance, and reliable data distribution across cluster nodes. The data storage indirectly participates in the processing procedure as nodes typically process local data.

All operations, such as read and write, in the data storage are mediated by our data governance approach. The data storage is implemented using components such as *Hive*, *Trino*, and *HDFS*.

**Data Governance** The framework of controls, including processes, policies, standards, and metrics, ensures efficient and secure data access. It implements access control policies during data ingestion and analytics using the Extract, Transform, Load (ETL) mechanism. The key components Ranger and Atlas enable effective data access controls and support the data governance framework. Ranger provides a centralized platform for defining and enforcing access policies on resources and tags, while Atlas implements annotations and data lineage. Policies in Ranger mediate access to resources, allowing for data governance flexibility. The result of a policy execution is a data transformation that addresses specific needs emerging in the considered scenario. We note that data transformations refer to Hive transformations. Hive, in fact, provides several manipulation functions that the user can extend to define her own User-Defined Functions (UDFs).

All functions, both those provided by the system and UDFs, can be used as custom masking option in Ranger.

**Data analytics** The set of components responsible for implementing the data analytics pipeline built on the data available on storage. It implements the big data pipelines analyzing data in the data storage and relies on

components *Spark* and *Airflow* as follows.

- *Spark* is an advanced, unified analytics engine for large-scale data processing. The core idea of Spark is to provide an expressive computing system (not limited to the Hadoop MapReduce model) by exploiting in-memory processing of cached data to avoid saving intermediate results to disk and caching data for repetitive queries. Apache Spark framework enables four different programming languages for map-reduce programming (i.e., Java, Scala, Python, R). We considered Scala mainly thanks to the greater documentation support and community.
- *Airflow* is an orchestrator of processing pipelines that aims to provide scheduling and monitoring capabilities. Airflow pipelines are implemented in Python for better dynamic pipeline generation. The pipeline is modeled as a Direct Acyclic Graph (DAG) of tasks that can be Spark jobs.

We recall that each access to data by big data pipelines is mediated by our data governance approach.

## 6 Experimental Evaluation

We experimentally evaluated our approach in the smart city reference scenario presented in Section 2.2. In the following, we first present the considered dataset and our experimental settings (Section 6.1); then, we present a concrete execution of our data governance approach in terms of access control policy enforcement based on spatial transformations (Section 6.2); we further explain the training of 5 anomaly detection models on 5 different spatial views of the considered dataset and the results of their execution at inference time (Section

6.3); finally, we discuss how policy enforcement affects the privacy and quality of the overall analytics (Section 6.4). We note that the 5 different detection models represent the 5 different configurations of our anomaly detector and, in turn, the multi-tenancy proper of our reference scenario.

## 6.1 Experimental Settings

We tested our methodology in a smart city scenario using a data set taken from open data on local public transport in Oslo. The dataset counts 9.88M rows and 25 columns (attributes). Each row contains information about the location of a single bus/tram and other information such as origin, destination, delay, and malfunctioning. All attributes in the dataset are detailed in Table 1.

An anomaly detector was implemented to spot abnormal situations, that is, any situation diverging from the *normal* behavior. The goal of the detector is to retrieve early indicators of critical events such as accidents, unauthorized protests, or terrorist attacks.

In the above settings, we considered that the traffic anomaly detector could be used by three categories of users (i.e., policeman, air quality control officer, citizen) under two different environment conditions (i.e., *normal*, *emergency*), leading to 5 different access control policies. These policies mediate access to data by means of spatial transformations that limit the view of the detector on the data.

Our experiments were run on a single node virtual machine running Ubuntu 20.04.2 LTS, installing the big data engine in Section 5. The virtual machine was equipped with Intel Xeon E5 2620 - 2.095GHZ CPU and 64 GB of RAM.

## 6.2 Data Transformation: Aggregation and Filtering

The data transformation process consists of a preparation and a filtering phase, working as follows.

**Aggregation.** Data aggregation is a fundamental step for anomaly detection on time series and properly defines the concept of *time-unit-anomaly*. Our experiments represented and collected data on movable points (buses and trams) from a “*fixed point of view*” (i.e., a specific area of the city), and aggregated them in a 5-minute time window representing our time-unit-anomaly.

Starting from the vehicles traffic dataset in Table 1, data were prepared for analytics according to a time aggregation executed at ingestion time and producing the dataset in Table 2. Collected data were aggregated in 5-minute time windows considering data coming from  $n$  distinct locations.  $k < n$  was then selected, where  $k$  represents the number of spatial areas (spatial clusters) for spatial aggregation. Each spatial area is represented in terms of its cluster centroid, that is, a set of GPS (latitude and longitude) coordinates. Obviously, changing the value of  $k$  also changes the number of zones and centroids considered, and of the level of spacial precision: a larger  $k$  gives larger granularity and thus more precise information, whereas a smaller  $k$  leads to smaller granularity and thus less precise information. Our experiments used five different predictive models based on *k-means* trained according to five different aggregated datasets with  $k \in \{5, 10, 25, 50, 100\}$ .

Properly modeling the spatial dimensions of the data considering spatial autocorrelation phenomena<sup>25,26,27</sup> emphasizes possible agreements or discrepancies related to the different sensors measurements located in different positions. Therefore, according to the considered clustering model, we collected different aggregated data for training the final predictive models.

**Filtering.** We exploited the values of  $k$  to define the access control policies per-



forming spatial filtering on data according to the role (i.e., policeman, air quality control officer, citizen) and the conditions of the environment (*normal*, *emergency*). Intuitively, a user with more permissions can use a larger  $k$ , whether a user with fewer privileges should use a smaller  $k$ . Filtering on the aggregated and tagged data was performed at ingestion time by changing  $k$  as specified in the following policies.

**Policy 1:**  $\langle \text{Citizen}, \text{TrafficData}, \text{Read}, \text{Normal}, \text{AggregationClusters}=5 \rangle$ , with  $\text{Citizen} \equiv \langle \text{user}, \{(\text{department}=\text{"Citizen"})\} \rangle$  and  $\text{TrafficData} \equiv \langle \text{file}, \{(\text{tag}=\text{"traffic"})\} \rangle$

This policy specifies that *citizens* can access low-accurate data in a *normal* situation. Data quantization is very low (5 centroids), thus the retrieved data show only approximately the anomalous area.

**Policy 2:**  $\langle \text{Citizen}, \text{TrafficData}, \text{Read}, \text{Emergency}, \text{AggregationClusters}=10 \rangle$ , where  $\text{CitizenUser} \equiv \langle \text{user}, \{(\text{department}=\text{"Citizen"})\} \rangle$  and  $\text{TrafficData} \equiv \langle \text{file}, \{(\text{tag}=\text{"traffic"})\} \rangle$

With this policy, during an *emergency*, *citizens* can access medium-low-accurate data. Data quantization is low (10 centroids). Retrieved data allows to see the anomalous area.

**Policy 3:**  $\langle \text{AirQualityUser}, \text{TrafficData}, \text{Read}, \text{Normal}, \text{AggregationClusters}=25 \rangle$ , where

$\text{AirQualityUser} \equiv \langle \text{user}, \{(\text{department}=\text{"AirQuality"})\} \rangle$  and  $\text{TrafficData} \equiv \langle \text{file}, \{(\text{tag}=\text{"traffic"})\} \rangle$

In this case, the policy states that *air quality control officers* can access medium-accurate data both in *normal* and *emergency* situations. Data quantization is medium (25 centroids). Retrieved data are still useful to outline the zone where an anomaly occurred.

**Policy 4:**  $\langle \text{Policeman}, \text{TrafficData}, \text{Read}, \text{Normal}, \text{AggregationClusters}=50 \rangle$ , with  $\text{Policeman} \equiv \langle \text{user}, \{(\text{department}=\text{"Police"})\} \rangle$  and  $\text{TrafficData} \equiv \langle \text{file}, \{(\text{tag}=\text{"traffic"})\} \rangle$

The policy states that *policemen* can access accurate data in a *normal* situation. Data quantization is high (50 centroids).

**Policy 5:**  $\langle \text{Policeman}, \text{TrafficData}, \text{Read}, \text{Emergency}, \text{AggregationClusters}=100 \rangle$ , where

$\text{Policeman} \equiv \langle \text{user}, \{(\text{department}=\text{"Police"})\} \rangle$  and  $\text{TrafficData} \equiv \langle \text{file}, \{(\text{tag}=\text{"traffic"})\} \rangle$

During an *emergency*, *policemen* can access the most accurate data. Data quantization is very high (100 centroids) resulting in a more accurate model guaranteeing a more precise identification of the anomaly location.

Policy enforcement led to the generation of the 5 different datasets in Figure 6.

### 6.3 Data Analytics: Model Training and Inference

Our experiments used 5 time series data collected from March, 1st, 2022 at 17:40 to October, 18th, 2022 at 14:45. The training set included time series from March, 1st, 2022 at 17:40 to the end of September 2022; the test set included the remaining 18 days of October. The testing time series were randomly perturbed by considering for each test instance 100 times the real value of average delay, average in panic, and the average in congestion buses. 1% of the testing instances were randomly selected and perturbed. The output files of the experiments are available online.<sup>2</sup>

Our training process differs from traditional processes where a single model is trained on the original dataset and the inference results are aggregated and filtered according to users' privileges. It rather builds on 5 different models trained on the 5 different datasets generated at ingestion time, which natively address the access control requirements. Each model (M1–M5 in Table 3) corresponds to the enforcement of a policy (Policy 1–Policy 5 in Section 6.2). In particular, our analytics pipeline performs a training phase on the 5 subset of data obtained after spatial filtering (representing normal situations) and generates 5

<sup>2</sup>Output files are available at: <http://www.di.uniba.it/~mignone/materials/MBDAaaS/anomdet.7z>

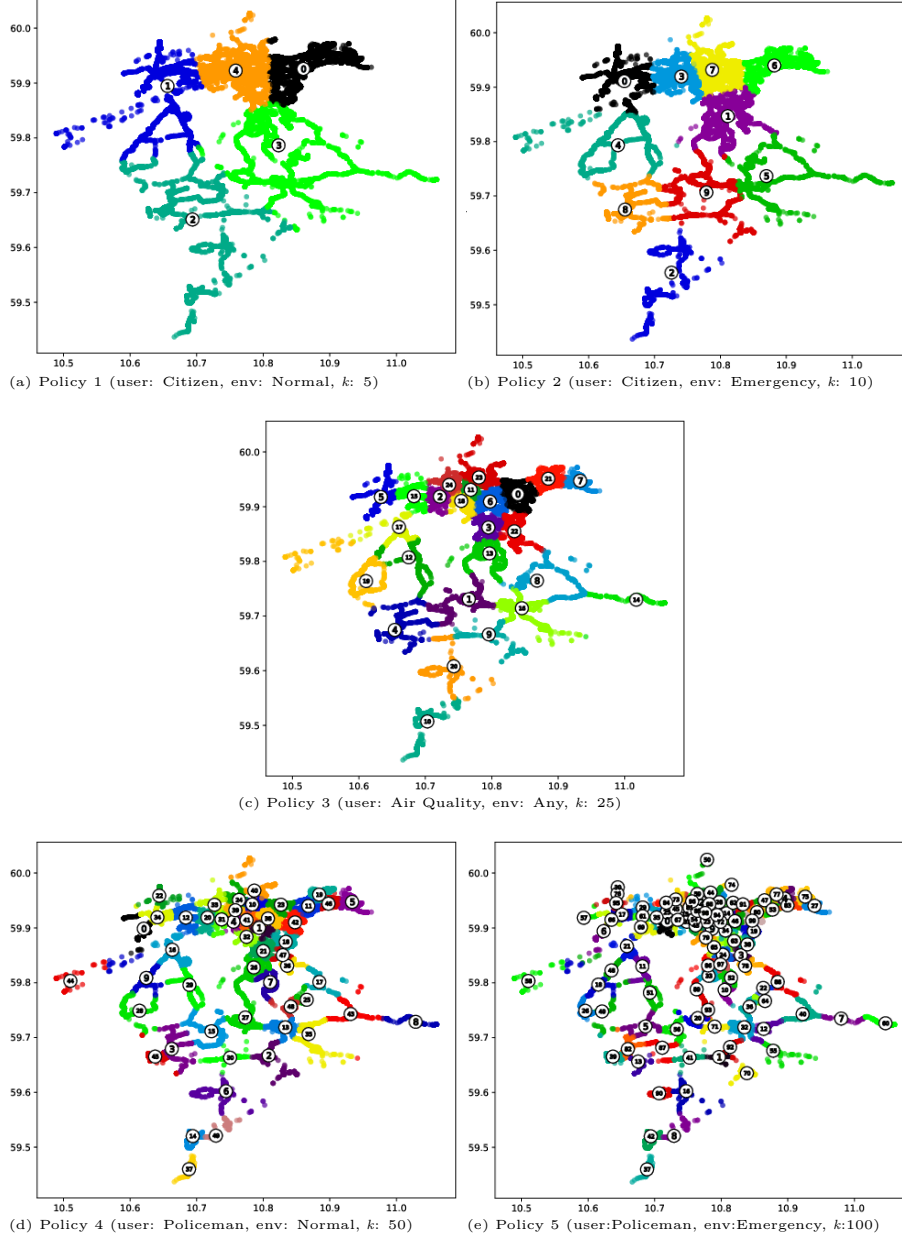


Figure 6: Results of the enforcement of policies in Section 6.2. Spatial transformation (number  $k$  of clusters) depends on the user role (*user*) and the status of the environment (*env*).

predictive models. At this point, each time a specific user decides to perform monitoring through the anomaly detector, it uses the model corresponding to his/her role. In this way, the model natively addresses data protection because it is built with data aggregated according to the access control policies based on the user’s privileges and only responses that can be accessed by the user requesting the analytics are used. In other words, the system presents as much sensitive data as possible and the detected anomalies at the level of detail granted by the user’s privileges, thus balancing operativity and GDPR compliance.

We note that we automatically configured the threshold factor  $t$  in Section 4.1 by using the trained GHSOM for predictions on the training set to achieve a given percentage  $fp\%$  of false positives, that is, zero false alarms. This is coherent with the assumption that the considered training set represents the set of historical normal cases<sup>11</sup>. Although this approach solves the problem of manually defining the appropriate value for  $t$ , it still requires defining  $fp\%$ , which is, anyway, much easier to define by the end-user than  $t$ . In our experiments, the best  $fp\%$  value is identified by maximising the f1-score in an internal cross-validation for  $fp\% \in \{0, 0.01, 0.1\}$ .

Table 3 shows the results in terms of weighted precision, recall, and f1-score w.r.t. the support of the classes normal and anomaly. f1-score is particularly suitable for the imbalanced task at hand due to the rare anomalous items and abundant normal cases. We considered the different predictive models  $M_1$ – $M_5$  generated according to the number of data clusters (5, 10, 25, 50, 100, resp.) imposed by Policy 1–Policy 5 transformations in Section 6.2. The results in Table 3 show that the number of clusters considered has a very low impact on the predictive performance (precision between 0.981 and 0.998, recall between 0.983 and 0.997), while it has a high impact on the amount of retrieved information.  $M_1$  and  $M_2$  being built on 5 and 10 clusters, resp., provide a rough idea of the

area affected by an anomaly (e.g., north-south-east-west-downtown);  $M_3$  being built on 25 clusters provides a more detailed idea of the area affected by an anomaly (e.g., district);  $M_4$  and  $M_5$  being built on 50 and 100 clusters, resp., provide a very detailed idea of the area affected by an anomaly (e.g., street, precise location).

Finally, we observed the results of the ablation analysis aimed to identify the best value for  $fp\%$  and, therefore, of  $t$ . The experiments showed that, the more the clusters, the data sensitivity, and the data volume, the higher the value of  $fp\%$  that maximizes the f1-score. Indeed, with  $k = 5$ ,  $k = 10$  and  $k = 25$ , the best results are obtained with a threshold  $t$  that is identified by assuming no false positives on the training set. On the contrary, with  $k = 50$  and  $k = 100$ , it is necessary to admit a small percentage of false positives on the training set to avoid too conservative models that do not detect abnormal cases properly. This behavior was somehow expected since a larger number of clusters requires the algorithm to consider cluster boundaries that are not clearly identifiable.

## 6.4 Discussion

Our main research objective focused on how to achieve a better balance between protection and quality in cloud-based and multi-tenant big data pipelines. That is, we aimed to ensure that the sensitive data used in the pipeline was protected even in case they are processed by users with different privileges, still producing high quality results. Our key findings are as follows.

**Full Privacy Scenario.** It considers models  $M_1$  and  $M_2$  to achieve a high degree of anonymisation at the expense of a substantial loss of quality in the data. What is important to note is that aggregation, performed before the training phase, actually changes the work and results of anomaly detection. In the case study, the system monitors areas of different sizes for

each level of aggregation (Figures 6(a) and 6(b)). The citizen’s view, for example, reduces the sensitivity of the model as small anomalies have less impact in a larger area. This behavior prevents the citizens from detecting anomalies he/she should not have access to and which are not of interest to him/her anyway.

**Full Quality Scenario.** It considers models M5 and M6 to greatly reduce anonymisation, while supporting a higher quality of data. As already discussed, aggregation carried out upstream changes the behavior of anomaly detection. In this case, the model is much more sensitive, making it possible to detect even point anomalies of greater interest to a security officer. We note that it is possible to take into account other signals (e.g., an emergency situation) to alter the granularity of the system, making it more flexible and adaptive.

**Full Post-Aggregation Scenario.** It considers the common approach with a single prediction model and a posteriori aggregation to anonymize the results. The model underlying such a system must be as granular as possible, avoiding the risk of removing information that is fundamental for the most privileged users. On the other hand, it is necessary to find an aggregation heuristic that is not a simple average over larger or smaller aggregation windows, so as to avoid an anomaly detected at higher granularities being visible also in those at lower granularities. This approach could also breach the rules imposed by the GDPR: although not all data are visible to the final end users, during the analytics they are managed by possibly unauthorized services or users.

## 7 Related Work

A growing number of companies and organizations use big data analytics tools to identify business opportunities and drive decision-making. Thus, the number of data security concerns has recently grown. It is commonly recognized that securing big data platforms takes a mix of traditional security tools, newly developed toolsets, and intelligent processes for managing and monitoring security throughout the entire data life-cycle<sup>28,29</sup> and this has consequently led to a new set of data security and privacy requirements<sup>30,7,31</sup>.

Most of the current solutions are based on Attribute Based Access Control (ABAC)<sup>14</sup> for its ability to define attribute-based policy rules, with attributes presented at run-time, coming from multiple sources. However, they suffer from limitations that range from focusing only on a subset of the requirements we highlighted in Section 2, or adding delay in updating and enforcing the security policies due to the increased complexity and computational overhead. Moreover, they are often tailored either to a very specific scenario or to a particular technology.

Some recent proposals, like<sup>9</sup> or<sup>32,33</sup>, propose a solution working only with the Apache Hadoop stack, leveraging on the native access control features of the platform, that still presents some limitations, such as the complexity of deployment and the consumption of resources. There are also database-oriented approaches, focusing on a particular database, hence on a particular type of analytics pipeline. For example,<sup>34</sup> presents an access control system for graph-based models, whereas<sup>35</sup> and<sup>8</sup> deal with NoSQL databases, such as MongoDB and HBase. All database-oriented approaches are based on query rewriting mechanisms to avoid leakage of sensitive resources at the cost of high computational complexity and low efficiency in enforcing security policies. Scenario-specific approaches often focus on specific big data scenarios, such as the case

of federated cloud and edge Microservices<sup>36</sup> or Internet of Things (IoT) use cases<sup>37</sup>. The challenge mainly addressed by these works is the management of access control decisions in a scenario with multiple stakeholders in continuously evolving federations.

Finding unusual patterns or events in data that change over time and space, like urban traffic patterns, is a common problem faced in various fields. These anomalies, such as traffic congestion or large crowds, can be difficult to detect due to their rarity and the fact that their definition varies based on location and time. Current methods for identifying anomalies in this type of data often fail to consider the relationships between different points in space and time and the specific characteristics of the anomalies themselves<sup>38</sup>.

For instance, in<sup>39</sup>, the authors proposed a decomposing approach to detect anomalies of different types, such as abnormal pedestrian flows or traffic accidents with varying locations and times. Specifically, they distinguish between the normal component, i.e., urban dynamics decided by spatiotemporal features, and the abnormal component that is caused by anomalous events. In<sup>40</sup>, the authors proposed a framework that provides support for the exploration of multidimensional road traffic data through visual analytics. The anomaly detection method is based on a traditional SOM used for clustering. The number of clusters is optimized through Silhouette cluster analysis. This approach is inspired by previous work presented in<sup>41</sup> which used a traditional approach for anomaly detection based on clustering algorithms, SOMs, and Gaussian Mixture Models (GMMs). However these fail to consider the relationships between different points in space and time and the specific characteristics of the anomalies themselves.

In<sup>42</sup>, the authors proposed a survey on research in urban anomaly analytics discussing various types of anomalies in the analyzed contexts such as urban



anomalies, traffic anomalies, unexpected crowds, environment anomalies, and individual anomalies. Furthermore, the authors emphasized that one of the open challenges is posed by urban data that is full of noise for precise predictions. The consequence is that anomalous events can easily be less represented than noise. The anomaly detection method considered in this paper is aimed at resolving the open problems presented in<sup>42</sup> by taking into account additional information coming from the spatial and temporal information. Such additional information allows the system to exploit spatial proximity information or temporal recurrence/periodicity in identifying truly anomalous cases.

## 8 Conclusions

We presented a novel data engine solution based on advanced access control-based ETL, supporting fine-grained data management and protection. We detailed how we implemented it within a complete Apache-based Big Data engine. Adding access control during the data extraction process offers several benefits with respect to traditional solutions segregating users within the data lake and building different extraction processes. First of all, it enhances data governance by providing more control over data, reducing the risk of unauthorized access or misuse of information, and protecting sensitive data from unauthorized disclosure. Moreover, it enables data customization, allowing users to have personalized views or subsets of data relevant to their needs, improving decision-making. We evaluated the impact in terms of the performance of our solution in a case study in the smart city domain. Specifically, we considered the task of identifying anomalies in geo-referenced and time-stamped data, continuously produced by sensors. The experimental evaluation shows that the identified anomalies, built on data manipulated according to access control policies, provide different perspectives of the anomalies, also at different levels of

granularity, revealing only relevant information for the specific role/user. For future work, we plan to extend the study to other machine learning tasks for data generated by sensors, other than anomaly detection, such as time series classification and regression.

## References

- [1] Qureshi SR, Gupta A. Towards efficient Big Data and data analytics: A review. In: 2014 Conference on IT in Business, Industry and Government (CSIBIG); 2014. p. 1-6.
- [2] Samarati P. Protecting Respondents' Identities in Microdata Release. IEEE Transactions on Knowledge and Data Engineering. 2001.
- [3] Machanavajjhala A, Gehrke J, Kifer D, et al. L-diversity: privacy beyond k-anonymity. In: 22nd International Conference on Data Engineering (ICDE'06); 2006. p. 24-4.
- [4] Dwork C, McSherry F, Nissim K, et al. Calibrating Noise to Sensitivity in Private Data Analysis. Journal of Privacy and Confidentiality. 2017 May;7(3):17–51.
- [5] Basanta-Val P, Sánchez-Fernández L. Big-BOE: Fusing Spanish official gazette with big data technology. Big Data. 2018 Jun;6(2):124-38.
- [6] Basanta-Val P, Fernández-García N, Sánchez-Fernández L, et al. Patterns for Distributed Real-Time Stream Processing. IEEE Transactions on Parallel and Distributed Systems. 2017;28(11):3243-57.
- [7] Colombo P, Ferrari E. Access control technologies for Big Data management systems: literature review and future trends. Cybersecurity. 2019.

- [8] Huang L, Zhu Y, Wang X, et al. An Attribute-Based Fine-Grained Access Control Mechanism for HBase. In: Hartmann S, Küng J, Chakravarthy S, Anderst-Kotsis G, Tjoa AM, Khalil I, editors. Database and Expert Systems Applications. Cham: Springer International Publishing; 2019. p. 44-59.
- [9] Chen C, Elsayed MA, Zulkernine M. HBD-Authority: Streaming Access Control Model for Hadoop. In: 2020 IEEE 6th International Conference on Dependability in Sensor, Cloud and Big Data Systems and Application (DependSys); 2020. p. 16-25.
- [10] Anisetti M, Ardagna CA, Braghin C, et al. Dynamic and Scalable Enforcement of Access Control Policies for Big Data. In: Proceedings of the 13th International Conference on Management of Digital EcoSystems. MEDES '21. New York, NY, USA: Association for Computing Machinery; 2021. p. 71-78. Available from: <https://doi.org/10.1145/3444757.3485107>.
- [11] Malondkar A, Corizzo R, Kiringa I, et al. Spark-GHSOM: Growing Hierarchical Self-Organizing Map for large scale mixed attribute datasets. Information Sciences. 2018.
- [12] Ding W, Abdel-Basset M, Hawash H, et al. Explainability of artificial intelligence methods, applications and challenges: A comprehensive survey. Information Sciences. 2022.
- [13] Tjoa E, Guan C. A Survey on Explainable Artificial Intelligence (XAI): Toward Medical XAI. IEEE Transactions on Neural Networks and Learning Systems. 2021.
- [14] Hu VC, Ferraiolo D, Kuhn R, et al. Guide to Attribute Based Access Control (ABAC) Definition and Considerations. NIST special publication. 2014.

- [15] Yang Q, Liu Y, Chen T, et al. Federated Machine Learning: Concept and Applications. *ACM Trans Intell Syst Technol.* 2019 jan;10(2).
- [16] OASIS committee (Organization for the Advancement of Structured Information Standards). Rissanen E, editor. eXtensible Access Control Markup Language (XACML). OASIS Open; 2013. <http://www.oasis-open.org/committees/xacml/>.
- [17] Nieves M, Dempsey K, Pillitteri V. An Introduction to Information Security. NIST special publication. 2017 2017-06-22.
- [18] Goyal V, Pandey O, Sahai A, et al. Attribute-Based Encryption for Fine-Grained Access Control of Encrypted Data. In: *Proceedings of the 13th ACM Conference on Computer and Communications Security. CCS '06.* New York, NY, USA: Association for Computing Machinery; 2006. p. 89–98. Available from: <https://doi.org/10.1145/1180405.1180418>.
- [19] Boyen X, Waters B. Anonymous Hierarchical Identity-Based Encryption (Without Random Oracles). In: *Dwork C, editor. Advances in Cryptology - CRYPTO 2006.* Berlin, Heidelberg: Springer Berlin Heidelberg; 2006. p. 290-307.
- [20] Armknecht F, Boyd C, Carr C, et al. A Guide to Fully Homomorphic Encryption. *IACR Cryptology ePrint Archive.* 2015.
- [21] Kohonen T. Self-organized formation of topologically correct feature maps. *Biological Cybernetics.* 1982.
- [22] Kohonen T, Honkela T. Kohonen network. *Scholarpedia.* 2007.
- [23] Hsu CC. Generalizing self-organizing map for categorical data. *IEEE Transactions on Neural Networks.* 2006.

- [24] Kader GD, Perry M. Variability for Categorical Variables. *Journal of Statistics Education*. 2007.
- [25] Corizzo R, Ceci M, Pio G, et al. Spatially-Aware Autoencoders for Detecting Contextual Anomalies in Geo-Distributed Data. In: *Discovery Science: 24th International Conference, DS 2021, Halifax, NS, Canada, October 11–13, 2021, Proceedings*. Berlin, Heidelberg: Springer-Verlag; 2021. p. 461–471. Available from: [https://doi.org/10.1007/978-3-030-88942-5\\_36](https://doi.org/10.1007/978-3-030-88942-5_36).
- [26] Mignone P, Malerba D, Ceci M. Anomaly Detection for Public Transport and Air Pollution Analysis. In: *2022 IEEE International Conference on Big Data (Big Data)*; 2022. p. 2867-74.
- [27] Ceci M, Corizzo R, Japkowicz N, et al. ECHAD: Embedding-Based Change Detection From Multivariate Time Series in Smart Grids. *IEEE Access*. 2020.
- [28] Al-Badi A, Tarhini A, Khan AI. Exploring Big Data Governance Frameworks. *Procedia Computer Science*. 2018.
- [29] Ben Aissa MM, Sfaxi L, Robbana R. DECIDE: A New Decisional Big Data Methodology for a Better Data Governance. In: Themistocleous M, Papadaki M, Kamal MM, editors. *Information Systems*. Cham: Springer International Publishing; 2020. p. 63-78.
- [30] Anisetti M, Ardagna CA, Berto F. An assurance process for Big Data trustworthiness. *Future Generation Computer Systems*. 2023;146:34-46. Available from: <https://www.sciencedirect.com/science/article/pii/S0167739X23001371>.
- [31] Lv Z, Song H, Basanta-Val P, et al. Next-Generation Big Data Analytics:

- State of the Art, Challenges, and Future Research Topics. *IEEE Transactions on Industrial Informatics*. 2017;13(4):1891-9.
- [32] Gupta M, Patwa F, Sandhu R. An Attribute-Based Access Control Model for Secure Big Data Processing in Hadoop Ecosystem. In: *Proceedings of the Third ACM Workshop on Attribute-Based Access Control. ABAC'18*. New York, NY, USA: Association for Computing Machinery; 2018. p. 13–24. Available from: <https://doi.org/10.1145/3180457.3180463>.
- [33] Gupta M, Patwa F, Sandhu R. Object-Tagged RBAC Model for the Hadoop Ecosystem. In: Livraga G, Zhu S, editors. *Data and Applications Security and Privacy XXXI*. Cham: Springer International Publishing; 2017. p. 63–81.
- [34] Chabin J, Ciferri CDA, Halfeld-Ferrari M, et al. Role-Based Access Control on Graph Databases. In: *SOFSEM 2021: Theory and Practice of Computer Science: 47th International Conference on Current Trends in Theory and Practice of Computer Science, SOFSEM 2021, Bolzano-Bozen, Italy, January 25–29, 2021, Proceedings*. Berlin, Heidelberg: Springer-Verlag; 2021. p. 519–534. Available from: [https://doi.org/10.1007/978-3-030-67731-2\\_38](https://doi.org/10.1007/978-3-030-67731-2_38).
- [35] Gupta E, Sural S, Vaidya J, et al. Enabling Attribute-based Access Control in NoSQL Databases. *IEEE Transactions on Emerging Topics in Computing*. 2022.
- [36] Preuveneers D, Joosen W. Towards Multi-party Policy-based Access Control in Federations of Cloud and Edge Microservices. In: *2019 IEEE European Symposium on Security and Privacy Workshops (EuroS&PW)*; 2019. p. 29–38.

- [37] de Matos E, Tiburski RT, Amaral LA, et al. Providing Context-Aware Security for IoT Environments Through Context Sharing Feature. In: 2018 17th IEEE International Conference On Trust, Security And Privacy In Computing And Communications/ 12th IEEE International Conference On Big Data Science And Engineering (TrustCom/BigDataSE); 2018. p. 1711-5.
- [38] Sofuoglu SE, Aveyente S. GLOSS: Tensor-based anomaly detection in spatiotemporal urban traffic data. Signal Processing. 2022. Available from: <https://www.sciencedirect.com/science/article/pii/S0165168421004072>.
- [39] Zhang M, Li T, Shi H, et al. A Decomposition Approach for Urban Anomaly Detection Across Spatiotemporal Data. In: Proceedings of the Twenty-Eighth International Joint Conference on Artificial Intelligence, IJCAI-19. International Joint Conferences on Artificial Intelligence Organization; 2019. p. 6043-9. Available from: <https://doi.org/10.24963/ijcai.2019/837>.
- [40] Riveiro M, Lebram M, Elmer M. Anomaly Detection for Road Traffic: A Visual Analytics Framework. IEEE Transactions on Intelligent Transportation Systems. 2017.
- [41] Kraiman JB, Arouh SL, Webb ML. Automated anomaly detection processor. In: Sisti AF, Trevisani DA, editors. Enabling Technologies for Simulation Science VI. vol. 4716. International Society for Optics and Photonics. SPIE; 2002. p. 128 137. Available from: <https://doi.org/10.1117/12.474940>.
- [42] Zhang M, Li T, Yu Y, et al. Urban Anomaly Analytics: Description, Detection, and Prediction. IEEE Transactions on Big Data. 2022.

Table 1: Attributes and related descriptions of the dataset used in the experiments.

| Attribute                           | Description                                                                                                 |
|-------------------------------------|-------------------------------------------------------------------------------------------------------------|
| <i>DateTime</i>                     | Timestamp specifying when the position/status was recorded/updated                                          |
| <i>LinkDistance</i>                 | Distance in meters between the previous stop (or current, if located at stop) and the next stop             |
| <i>Percentage</i>                   | How much of the total distance (percentage) that has been traversed at the time of the message              |
| <i>LineRef</i>                      | Reference to the line in question                                                                           |
| <i>DirectionRef</i>                 | Reference to the direction in question                                                                      |
| <i>PublishedLineName</i>            | Name describing the line in question                                                                        |
| <i>OriginRef</i>                    | Reference to the Origin in question                                                                         |
| <i>OriginName</i>                   | Name describing the origin of the departure                                                                 |
| <i>DestinationRef</i>               | Reference to the destination in question                                                                    |
| <i>DestinationName</i>              | Name describing the destination of the departure                                                            |
| <i>OriginAimedDepartureTime</i>     | Origin aimed departure time                                                                                 |
| <i>DestinationAimedArrivalTime</i>  | Destination aimed arrival time                                                                              |
| <i>VehicleRef</i>                   | Reference to the vehicle in question                                                                        |
| <i>Delay</i>                        | Delay-time, defined as “PT0S” (0 seconds) when there are no delays                                          |
| <i>HeadwayService</i>               | Field defining whether the service is a headway service                                                     |
| <i>InCongestion</i>                 | Field defining whether the traffic is in congestion or other circumstances which may lead to further delays |
| <i>InPanic</i>                      | Boolean field, indicating that the driver reported an “out of order” status                                 |
| <i>GPS</i> (Latitude and Longitude) | Position of the public vehicle, according to the timestamp recorded                                         |
| <i>StopPointRef</i>                 | Reference to the stop point in question                                                                     |
| <i>VisitNumber</i>                  | Number of the stop point in question                                                                        |
| <i>StopPointName</i>                | Name of the stop point in question                                                                          |
| <i>VehicleAtStop</i>                | Field defining whether the vehicle is at the stop                                                           |
| <i>DestinationDisplay</i>           | Name of the next destination                                                                                |



Table 2: Aggregated dataset.

| Attribute                          | Description                                                                                                             |
|------------------------------------|-------------------------------------------------------------------------------------------------------------------------|
| <i>Date</i>                        | The timestamp of the current measurement/instance                                                                       |
| <i>ClusterLatitude</i>             | The latitude of the cluster centroid from where aggregate                                                               |
| <i>ClusterLongitude</i>            | The longitude of the cluster centroid from where aggregate                                                              |
| <i>Delay</i>                       | Average of the busses                                                                                                   |
| <i>Percentage</i>                  | Average total distance that has been traversed                                                                          |
| <i>InPanic</i>                     | Average busses in panic mode                                                                                            |
| <i>InCongestion</i>                | Average busses in congestion                                                                                            |
| <i>DestinationAimedArrivalTime</i> | Average Destination aimed departure time                                                                                |
| <i>OriginAimedDepartureTime</i>    | Average Origin aimed departure time                                                                                     |
| <i>HeadwayService_False</i>        | Count of instances for non-headway services                                                                             |
| <i>Anomaly</i>                     | 0,1 indicating if the instance represents a normal case<br>0 or an anomaly 1 (for quantitative evaluation purpose only) |

Table 3: Results in terms of precision, recall, and f1-score. We also report the training time (in seconds), the training set cardinality, as well as the value for  $fp\%$  that maximizes the f1-score.

| model            | prec  | rec   | f1-score | fp%  | train time | data card |
|------------------|-------|-------|----------|------|------------|-----------|
| $M_1$ (5_clus)   | 0.995 | 0.993 | 0.994    | 0    | 140        | 198312    |
| $M_2$ (10_clus)  | 0.998 | 0.997 | 0.997    | 0    | 362        | 399677    |
| $M_3$ (25_clus)  | 0.990 | 0.990 | 0.986    | 0    | 1997       | 910724    |
| $M_4$ (50_clus)  | 0.981 | 0.983 | 0.982    | 0.1  | 1488       | 1577192   |
| $M_5$ (100_clus) | 0.982 | 0.990 | 0.985    | 0.01 | 7888       | 2546608   |