

Anomaly Detection: A Survey

VARUN CHANDOLA, ARINDAM BANERJEE, and VIPIN KUMAR

University of Minnesota

15

Anomaly detection is an important problem that has been researched within diverse research areas and application domains. Many anomaly detection techniques have been specifically developed for certain application domains, while others are more generic. This survey tries to provide a structured and comprehensive overview of the research on anomaly detection. We have grouped existing techniques into different categories based on the underlying approach adopted by each technique. For each category we have identified key assumptions, which are used by the techniques to differentiate between normal and anomalous behavior. When applying a given technique to a particular domain, these assumptions can be used as guidelines to assess the effectiveness of the technique in that domain. For each category, we provide a basic anomaly detection technique, and then show how the different existing techniques in that category are variants of the basic technique. This template provides an easier and more succinct understanding of the techniques belonging to each category. Further, for each category, we identify the advantages and disadvantages of the techniques in that category. We also provide a discussion on the computational complexity of the techniques since it is an important issue in real application domains. We hope that this survey will provide a better understanding of the different directions in which research has been done on this topic, and how techniques developed in one area can be applied in domains for which they were not intended to begin with.

Categories and Subject Descriptors: H.2.8 [Database Management]: Database Applications—*Data mining*

General Terms: Algorithms

Additional Key Words and Phrases: Anomaly detection, outlier detection

ACM Reference Format:

Chandola, V., Banerjee, A., and Kumar, V. 2009. Anomaly detection: A survey. *ACM Comput. Surv.* 41, 3, Article 15 (July 2009), 58 pages. DOI = 10.1145/1541880.1541882
<http://doi.acm.org/10.1145/1541880.1541882>

1. INTRODUCTION

Anomaly detection refers to the problem of finding patterns in data that do not conform to expected behavior. These nonconforming patterns are often referred to as anomalies, outliers, discordant observations, exceptions, aberrations, surprises, peculiarities, or

This work was supported by NASA under award NNX08AC36A, NSF grant number CNS-0551551, NSF ITR Grant ACI-0325949, NSF IIS-0713227, and NSF Grant IIS-0308264. Access to computing facilities was provided by the Digital Technology Consortium.

Author's address: V. Chandola, Computer Science Department, University of Minnesota, 200 Union St., Minneapolis, MN 55414.

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies show this notice on the first page or initial screen of a display along with the full citation. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers, to redistribute to lists, or to use any component of this work in other works requires prior specific permission and/or a fee. Permissions may be requested from Publications Dept., ACM, Inc., 2 Penn Plaza, Suite 701, New York, NY 10121-0701 USA, fax +1 (212) 869-0481, or permissions@acm.org.

©2009 ACM 0360-0300/2009/07-ART15 \$10.00

DOI 10.1145/1541880.1541882 <http://doi.acm.org/10.1145/1541880.1541882>

ACM Computing Surveys, Vol. 41, No. 3, Article 15, Publication date: July 2009.

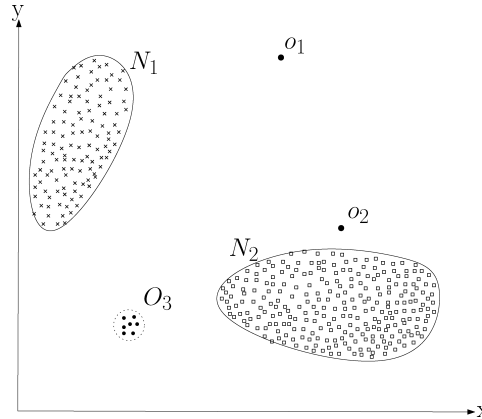


Fig. 1. A simple example of anomalies in a two-dimensional data set.

contaminants in different application domains. Of these, anomalies and outliers are two terms used most commonly in the context of anomaly detection; sometimes interchangeably. Anomaly detection finds extensive use in a wide variety of applications such as fraud detection for credit cards, insurance, or health care, intrusion detection for cyber-security, fault detection in safety critical systems, and military surveillance for enemy activities.

The importance of anomaly detection is due to the fact that anomalies in data translate to significant, and often critical, actionable information in a wide variety of application domains. For example, an anomalous traffic pattern in a computer network could mean that a hacked computer is sending out sensitive data to an unauthorized destination [Kumar 2005]. An anomalous MRI image may indicate the presence of malignant tumors [Spence et al. 2001]. Anomalies in credit card transaction data could indicate credit card or identity theft [Aleskerov et al. 1997], or anomalous readings from a space craft sensor could signify a fault in some component of the space craft [Fujimaki et al. 2005].

Detecting outliers or anomalies in data has been studied in the statistics community as early as the 19th century [Edgeworth 1887]. Over time, a variety of anomaly detection techniques have been developed in several research communities. Many of these techniques have been specifically developed for certain application domains, while others are more generic.

This survey tries to provide a structured and comprehensive overview of the research on anomaly detection. We hope that it facilitates a better understanding of the different directions in which research has been done on this topic, and how techniques developed in one area can be applied in domains for which they were not intended to begin with.

1.1. What are Anomalies?

Anomalies are patterns in data that do not conform to a well defined notion of normal behavior. Figure 1 illustrates anomalies in a simple two-dimensional data set. The data has two normal regions, N_1 and N_2 , since most observations lie in these two regions. Points that are sufficiently far away from these regions, for example, points o_1 and o_2 , and points in region O_3 , are anomalies.

Anomalies might be induced in the data for a variety of reasons, such as malicious activity, for example, credit card fraud, cyber-intrusion, terrorist activity or breakdown of a system, but all of the reasons have the common characteristic that they are

interesting to the analyst. The interestingness or real life relevance of anomalies is a key feature of anomaly detection.

Anomaly detection is related to, but distinct from *noise removal* [Teng et al. 1990] and *noise accommodation* [Rousseeuw and Leroy 1987], both of which deal with unwanted *noise* in the data. Noise can be defined as a phenomenon in data that is not of interest to the analyst, but acts as a hindrance to data analysis. Noise removal is driven by the need to remove the unwanted objects before any data analysis is performed. Noise accommodation refers to immunizing a statistical model estimation against anomalous observations [Huber 1974].

Another topic related to anomaly detection is *novelty detection* [Markou and Singh 2003a, 2003b; Saunders and Gero 2000], which aims at detecting previously unobserved (*emergent, novel*) patterns in the data, for example, a new topic of discussion in a news group. The distinction between novel patterns and anomalies is that the novel patterns are typically incorporated into the normal model after being detected.

It should be noted that solutions for these related problems are often used for anomaly detection and vice versa, and hence are discussed in this review as well.

1.2. Challenges

At an abstract level, an anomaly is defined as a pattern that does not conform to expected normal behavior. A straightforward anomaly detection approach, therefore, is to define a region representing normal behavior and declare any observation in the data that does not belong to this normal region as an anomaly. But several factors make this apparently simple approach very challenging:

- Defining a normal region that encompasses every possible normal behavior is very difficult. In addition, the boundary between normal and anomalous behavior is often not precise. Thus an anomalous observation that lies close to the boundary can actually be normal, and vice versa.
- When anomalies are the result of malicious actions, the malicious adversaries often adapt themselves to make the anomalous observations appear normal, thereby making the task of defining normal behavior more difficult.
- In many domains normal behavior keeps evolving and a current notion of normal behavior might not be sufficiently representative in the future.
- The exact notion of an anomaly is different for different application domains. For example, in the medical domain a small deviation from normal (e.g., fluctuations in body temperature) might be an anomaly, while similar deviation in the stock market domain (e.g., fluctuations in the value of a stock) might be considered as normal. Thus applying a technique developed in one domain to another, is not straightforward.
- Availability of labeled data for training/validation of models used by anomaly detection techniques is usually a major issue.
- Often the data contains noise that tends to be similar to the actual anomalies and hence is difficult to distinguish and remove.

Due to these challenges, the anomaly detection problem, in its most general form, is not easy to solve. In fact, most of the existing anomaly detection techniques solve a specific formulation of the problem. The formulation is induced by various factors such as the nature of the data, availability of labeled data, type of anomalies to be detected, and so on. Often, these factors are determined by the application domain in which the anomalies need to be detected. Researchers have adopted concepts from diverse disciplines such as statistics, machine learning, data mining, information theory, spectral

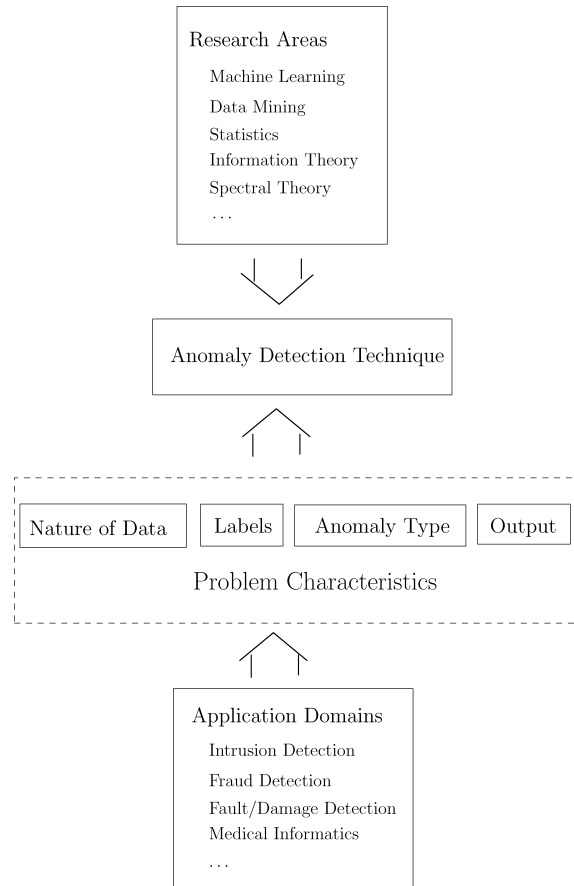


Fig. 2. Key components associated with an anomaly detection technique.

theory, and have applied them to specific problem formulations. Figure 2 shows the key components associated with any anomaly detection technique.

1.3. Related Work

Anomaly detection has been the topic of a number of surveys and review articles, as well as books. Hodge and Austin [2004] provide an extensive survey of anomaly detection techniques developed in machine learning and statistical domains. A broad review of anomaly detection techniques for numeric as well as symbolic data is presented by Agyemang et al. [2006]. An extensive review of novelty detection techniques using neural networks and statistical approaches has been presented in Markou and Singh [2003a] and Markou and Singh [2003b], respectively. Patcha and Park [2007] and Snyder [2001] present a survey of anomaly detection techniques used specifically for cyber-intrusion detection. A substantial amount of research on outlier detection has been done in statistics and has been reviewed in several books [Rousseeuw and Leroy 1987; Barnett and Lewis 1994; Hawkins 1980] as well as other survey articles [Beckman and Cook 1983; Bakar et al. 2006].

Table I shows the set of techniques and application domains covered by our survey and the various related survey articles.

Table I. Comparison of our Survey to Other Related Survey Articles. 1—Our Survey, 2—Hodge and Austin [2004], 3—Agyemang et al. [2006], 4—Markou and Singh [2003a], 5—Markou and Singh [2003b], 6—Patcha and Park [2007], 7—Beckman and Cook [1983], 8—Bakar et al. [2006]

		1	2	3	4	5	6	7	8
Techniques	Classification Based	✓	✓	✓	✓		✓		
	Clustering Based	✓	✓	✓			✓		
	Nearest Neighbor Based	✓	✓	✓			✓		✓
	Statistical	✓	✓	✓		✓	✓	✓	✓
	Information Theoretic	✓							
	Spectral	✓							
Applications	Cyber-Intrusion Detection	✓					✓		
	Fraud Detection	✓							
	Medical Anomaly Detection	✓							
	Industrial Damage Detection	✓							
	Image Processing	✓							
	Textual Anomaly Detection	✓							
	Sensor Networks	✓							

1.4. Our Contributions

This survey is an attempt to provide a structured and broad overview of extensive research on anomaly detection techniques spanning multiple research areas and application domains.

Most of the existing surveys on anomaly detection either focus on a particular application domain or on a single research area. Agyemang et al. [2006] and Hodge and Austin [2004] are two related works that group anomaly detection into multiple categories and discuss techniques under each category. This survey builds upon these two works by significantly expanding the discussion in several directions.

We add two more categories of anomaly detection techniques, information theoretic and spectral techniques, to the four categories discussed in Agyemang et al. [2006] and Hodge and Austin [2004]. For each of the six categories, we not only discuss the techniques, but also identify unique assumptions regarding the nature of anomalies made by the techniques in that category. These assumptions are critical for determining when the techniques in that category would be able to detect anomalies, and when they would fail. For each category, we provide a basic anomaly detection technique, and then show how the different existing techniques in that category are variants of the basic technique. This template provides an easier and more succinct understanding of the techniques belonging to each category. Further, for each category we identify the advantages and disadvantages of the techniques. We also provide a discussion of the computational complexity of the techniques since that is an important issue in real application domains.

While some of the existing surveys mention the different applications of anomaly detection, we provide a detailed discussion of the application domains where anomaly detection techniques have been used. For each domain we discuss the notion of an anomaly, the different aspects of the anomaly detection problem, and the challenges faced by the anomaly detection techniques. We also provide a list of techniques that have been applied in each application domain.

The existing surveys discuss anomaly detection techniques that detect the simplest form of anomalies. We distinguish simple anomalies from complex anomalies. The discussion of applications of anomaly detection reveals that for most application domains, the interesting anomalies are complex in nature, while most of the algorithmic research has focussed on simple anomalies.

1.5. Organization

This survey is organized into three parts and its structure closely follows Figure 2. In Section 2, we identify the various aspects that determine the formulation of the problem and highlight the richness and complexity associated with anomaly detection. We distinguish simple anomalies from complex anomalies and define two types of complex anomalies: contextual and collective anomalies. In Section 3, we briefly describe the different application domains to which anomaly detection has been applied. In subsequent sections we provide a categorization of anomaly detection techniques based on the research area to which they belong. A majority of the techniques can be categorized into classification-based (Section 4), nearest neighbor-based (Section 5), clustering-based (Section 6), and statistical techniques (Section 7). Some techniques belong to research areas such as information theory (Section 8), and spectral theory (Section 9). For each category of techniques we also discuss their computational complexity for training and testing phases. In Section 10 we discuss various contextual anomaly detection techniques. We present some discussion of the limitations and relative performance of various existing techniques in Section 11. Section 12 contains concluding remarks.

2. DIFFERENT ASPECTS OF AN ANOMALY DETECTION PROBLEM

This section identifies and discusses the different aspects of anomaly detection. As mentioned earlier, a specific formulation of the problem is determined by several different factors such as the nature of the input data, the availability or unavailability of labels as well as the constraints and requirements induced by the application domain. This section brings forth the richness in the problem domain and justifies the need for the broad spectrum of anomaly detection techniques.

2.1. Nature of Input Data

A key aspect of any anomaly detection technique is the nature of the input data. Input is generally a collection of data instances (also referred as *object*, *record*, *point*, *vector*, *pattern*, *event*, *case*, *sample*, *observation*, or *entity*) [Tan et al. 2005, Chapter 2]. Each data instance can be described using a set of attributes (also referred to as *variable*, *characteristic*, *feature*, *field*, or *dimension*). The attributes can be of different types such as *binary*, *categorical*, or *continuous*. Each data instance might consist of only one attribute (*univariate*) or multiple attributes (*multivariate*). In the case of multivariate data instances, all attributes might be of same type or might be a mixture of different data types.

The nature of attributes determines the applicability of anomaly detection techniques. For example, for statistical techniques different statistical models have to be used for continuous and categorical data. Similarly, for nearest-neighbor-based techniques, the nature of attributes would determine the distance measure to be used. Often, instead of the actual data, the pairwise distance between instances might be provided in the form of a distance or similarity matrix. In such cases, techniques that require original data instances are not applicable, for example, many statistical and classification-based techniques.

Input data can also be categorized based on the relationship present among data instances [Tan et al. 2005]. Most of the existing anomaly detection techniques deal with record data or point data, in which no relationship is assumed among the data instances.

In general, data instances can be related to each other. Some examples are *sequence data*, *spatial data*, and *graph data*. In sequence data, the data instances are linearly

ordered, for example, time-series data, genome sequences, and protein sequences. In *spatial data*, each data instance is related to its neighboring instances, for example, vehicular traffic data, and ecological data. When the spatial data has a temporal (sequential) component it is referred to as *spatio-temporal* data, for example, climate data. In *graph data*, data instances are represented as vertices in a graph and are connected to other vertices with edges. Later in this section we will discuss situations where such relationships among data instances become relevant for anomaly detection.

2.2. Type of Anomaly

An important aspect of an anomaly detection technique is the nature of the desired anomaly. Anomalies can be classified into following three categories:

2.2.1. Point Anomalies. If an individual data instance can be considered as anomalous with respect to the rest of data, then the instance is termed a point anomaly. This is the simplest type of anomaly and is the focus of majority of research on anomaly detection.

For example, in Figure 1, points o_1 and o_2 , as well as points in region O_3 , lie outside the boundary of the normal regions, and hence are point anomalies since they are different from normal data points.

As a real-life example, consider credit card fraud detection. Let the data set correspond to an individual's credit card transactions. For the sake of simplicity, let us assume that the data is defined using only one feature: *amount spent*. A transaction for which the amount spent is very high compared to the normal range of expenditure for that person will be a point anomaly.

2.2.2. Contextual Anomalies. If a data instance is anomalous in a specific context, but not otherwise, then it is termed a contextual anomaly (also referred to as *conditional anomaly* [Song et al. 2007]).

The notion of a context is induced by the structure in the data set and has to be specified as a part of the problem formulation. Each data instance is defined using the following two sets of attributes:

- (1) *Contextual attributes.* The contextual attributes are used to determine the context (or neighborhood) for that instance. For example, in spatial data sets, the longitude and latitude of a location are the contextual attributes. In time-series data, time is a contextual attribute that determines the position of an instance on the entire sequence.
- (2) *Behavioral attributes.* The behavioral attributes define the noncontextual characteristics of an instance. For example, in a spatial data set describing the average rainfall of the entire world, the amount of rainfall at any location is a behavioral attribute.

The anomalous behavior is determined using the values for the behavioral attributes within a specific context. A data instance might be a contextual anomaly in a given context, but an identical data instance (in terms of behavioral attributes) could be considered normal in a different context. This property is key in identifying contextual and behavioral attributes for a contextual anomaly detection technique.

Contextual anomalies have been most commonly explored in time-series data [Weigend et al. 1995; Salvador and Chan 2003] and spatial data [Kou et al. 2006; Shekhar et al. 2001]. Figure 3 shows one such example for a temperature time-series that shows the monthly temperature of an area over the last few years. A temperature

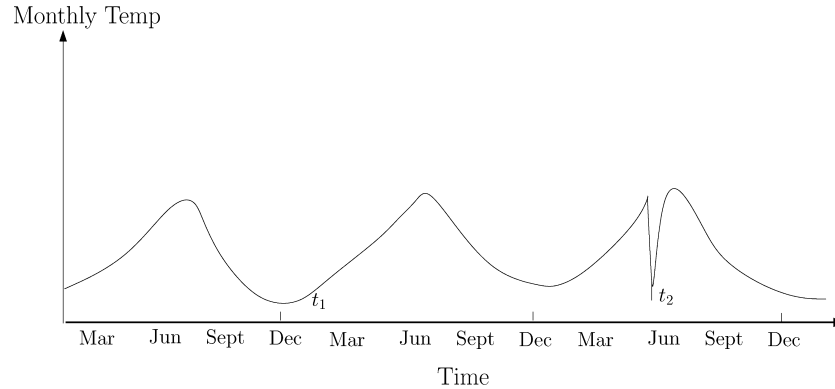


Fig. 3. Contextual anomaly t_2 in a temperature time-series. Note that the temperature at time t_1 is same as that at time t_2 but occurs in a different context and hence is not considered as an anomaly.

of 35°F might be normal during the winter (at time t_1) at that place, but the same value during the summer (at time t_2) would be an anomaly.

A similar example can be found in the credit card fraud detection domain. A contextual attribute in the credit card domain can be the time of purchase. Suppose an individual usually has a weekly shopping bill of \$100 except during the Christmas week, when it reaches \$1000. A new purchase of \$1000 in a week in July will be considered a contextual anomaly, since it does not conform to the normal behavior of the individual in the context of time even though the same amount spent during Christmas week will be considered normal.

The choice of applying a contextual anomaly detection technique is determined by the meaningfulness of the contextual anomalies in the target application domain. Another key factor is the availability of contextual attributes. In several cases defining a context is straightforward, and hence applying a contextual anomaly detection technique makes sense. In other cases, defining a context is not easy, making it difficult to apply such techniques.

2.2.3. Collective Anomalies. If a collection of related data instances is anomalous with respect to the entire data set, it is termed a collective anomaly. The individual data instances in a collective anomaly may not be anomalies by themselves, but their occurrence together as a collection is anomalous. Figure 4 is an example that shows a human electrocardiogram output [Goldberger et al. 2000]. The highlighted region denotes an anomaly because the same low value exists for an abnormally long time (corresponding to an *Atrial Premature Contraction*). Note that that low value by itself is not an anomaly.

As an another illustrative example, consider a sequence of actions occurring in a computer as shown below:

... http-web, buffer-overflow, http-web, http-web, smtp-mail, ftp, http-web, ssh, smtp-mail, http-web, **ssh, buffer-overflow, ftp**, http-web, ftp, smtp-mail, http-web ...

The highlighted sequence of events (**buffer-overflow, ssh, ftp**) correspond to a typical Web-based attack by a remote machine followed by copying of data from the host computer to a remote destination via *ftp*. It should be noted that this collection of events is an anomaly, but the individual events are not anomalies when they occur in other locations in the sequence.

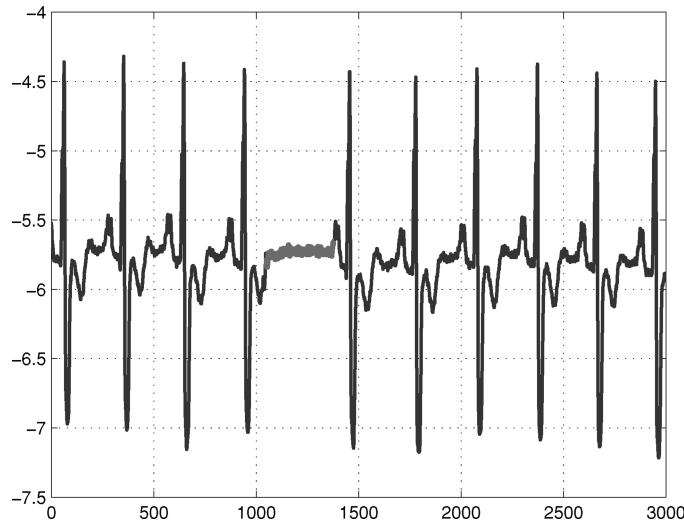


Fig. 4. Collective anomaly corresponding to an *Atrial Premature Contraction* in an human electrocardiogram output.

Collective anomalies have been explored for sequence data [Forrest et al. 1999; Sun et al. 2006], graph data [Noble and Cook 2003], and spatial data [Shekhar et al. 2001].

It should be noted that while point anomalies can occur in any data set, collective anomalies can occur only in data sets in which data instances are related. In contrast, occurrence of contextual anomalies depends on the availability of context attributes in the data. A point anomaly or a collective anomaly can also be a contextual anomaly if analyzed with respect to a context. Thus a point anomaly detection problem or collective anomaly detection problem can be transformed to a contextual anomaly detection problem by incorporating the context information.

The techniques used for detecting collective anomalies are very different than the point and contextual anomaly detection techniques, and require a separate detailed discussion. Hence we chose to not cover them in this survey. For a brief review of the research done in the field of collective anomaly detection, the reader is referred to an extended version of this survey [Chandola et al. 2007].

2.3. Data Labels

The labels associated with a data instance denote whether that instance is *normal* or *anomalous*.¹ It should be noted that obtaining labeled data that is accurate as well as representative of all types of behaviors, is often prohibitively expensive. Labeling is often done manually by a human expert and hence substantial effort is required to obtain the labeled training data set. Typically, getting a labeled set of anomalous data instances that covers all possible type of anomalous behavior is more difficult than getting labels for normal behavior. Moreover, the anomalous behavior is often dynamic in nature, for example, new types of anomalies might arise, for which there is no labeled training data. In certain cases, such as air traffic safety, anomalous instances would translate to catastrophic events, and hence are very rare.

Based on the extent to which the labels are available, anomaly detection techniques can operate in one of the following three modes:

¹Also referred to as normal and anomalous classes.

2.3.1. Supervised Anomaly Detection. Techniques trained in supervised mode assume the availability of a training data set that has labeled instances for normal as well as anomaly classes. A typical approach in such cases is to build a predictive model for normal vs. anomaly classes. Any unseen data instance is compared against the model to determine which class it belongs to. There are two major issues that arise in supervised anomaly detection. First, the anomalous instances are far fewer compared to the normal instances in the training data. Issues that arise due to imbalanced class distributions have been addressed in the data mining and machine learning literature [Joshi et al. 2001, 2002; Chawla et al. 2004; Phua et al. 2004; Weiss and Hirsh 1998; Vilalta and Ma 2002]. Second, obtaining accurate and representative labels, especially for the anomaly class is usually challenging. A number of techniques have been proposed that inject artificial anomalies into a normal data set to obtain a labeled training data set [Theiler and Cai 2003; Abe et al. 2006; Steinwart et al. 2005]. Other than these two issues, the supervised anomaly detection problem is similar to building predictive models. Hence we will not address this category of techniques in this survey.

2.3.2. Semisupervised Anomaly Detection. Techniques that operate in a semisupervised mode, assume that the training data has labeled instances only for the normal class. Since they do not require labels for the anomaly class, they are more widely applicable than supervised techniques. For example, in spacecraft fault detection [Fujimaki et al. 2005], an anomaly scenario would signify an accident, which is not easy to model. The typical approach used in such techniques is to build a model for the class corresponding to normal behavior, and use the model to identify anomalies in the test data.

A limited set of anomaly detection techniques exists that assumes availability of only the anomaly instances for training [Dasgupta and Nino 2000; Dasgupta and Majumdar 2002; Forrest et al. 1999]. Such techniques are not commonly used, primarily because it is difficult to obtain a training data set that covers every possible anomalous behavior that can occur in the data.

2.3.3. Unsupervised Anomaly Detection. Techniques that operate in unsupervised mode do not require training data, and thus are most widely applicable. The techniques in this category make the implicit assumption that normal instances are far more frequent than anomalies in the test data. If this assumption is not true then such techniques suffer from high false alarm rate.

Many semisupervised techniques can be adapted to operate in an unsupervised mode by using a sample of the unlabeled data set as training data. Such adaptation assumes that the test data contains very few anomalies and the model learned during training is robust to these few anomalies.

2.4. Output of Anomaly Detection

An important aspect for any anomaly detection technique is the manner in which the anomalies are reported. Typically, the outputs produced by anomaly detection techniques are one of the following two types:

2.4.1. Scores. Scoring techniques assign an anomaly score to each instance in the test data depending on the degree to which that instance is considered an anomaly. Thus the output of such techniques is a ranked list of anomalies. An analyst may choose to either analyze the top few anomalies or use a cutoff threshold to select the anomalies.

2.4.2. Labels. Techniques in this category assign a label (*normal* or *anomalous*) to each test instance.

```

open,  read,  mmap,  mmap,  open,  read,  mmap  ...
open,  mmap,  mmap,  read,  open,  close  ...
open,  close, open,  close, open,  mmap,  close  ...

```

Fig. 5. A sample data set comprised of three operating system call traces.

Scoring-based anomaly detection techniques allow the analyst to use a domain-specific threshold to select the most relevant anomalies. Techniques that provide binary labels to the test instances do not directly allow the analysts to make such a choice, though this can be controlled indirectly through parameter choices within each technique.

3. APPLICATIONS OF ANOMALY DETECTION

In this section we discuss several applications of anomaly detection. For each application domain we discuss the following four aspects:

- the notion of anomaly;
- nature of the data;
- challenges associated with detecting anomalies;
- existing anomaly detection techniques.

3.1. Intrusion Detection

Intrusion detection refers to detection of malicious activity (break-ins, penetrations, and other forms of computer abuse) in a computer related system [Phoha 2002]. These malicious activities or *intrusions* are interesting from a computer security perspective. An intrusion is different from the normal behavior of the system, and hence anomaly detection techniques are applicable in intrusion detection domain.

The key challenge for anomaly detection in this domain is the huge volume of data. The anomaly detection techniques need to be computationally efficient to handle these large sized inputs. Moreover the data typically comes in a streaming fashion, thereby requiring online analysis. Another issue that arises because of the large sized input is the false alarm rate. Since the data amounts to millions of data objects, a few percent of false alarms can make analysis overwhelming for an analyst. Labeled data corresponding to normal behavior is usually available, while labels for intrusions are not. Thus, semisupervised and unsupervised anomaly detection techniques are preferred in this domain.

Denning [1987] classifies intrusion detection systems into **host-based** and **network-based** intrusion detection systems.

3.1.1. Host-Based Intrusion Detection Systems. Such systems (also referred to as system call intrusion detection systems) deal with operating system call traces. The intrusions are in the form of anomalous subsequences (collective anomalies) of the traces. The anomalous subsequences translate to malicious programs, unauthorized behavior and policy violations. While all traces contain events belonging to the same alphabet, it is the co-occurrence of events that is the key factor in differentiating between normal and anomalous behavior.

The data is sequential in nature and the alphabet consists of individual system calls as shown in Figure 5. These calls could be generated by programs [Hofmeyr et al. 1998] or by users [Lane and Brodley 1999]. The alphabet is usually large (183 system calls for SunOS 4.1x Operating System). Different programs execute these system calls in

Table II. Examples of Anomaly Detection Techniques Used for Host-Based Intrusion Detection

Technique Used	Section	References
Statistical Profiling using Histograms	Section 7.2.1	Forrest et al. [1996a, 1999, 2004, 1994, 1999], Hofmeyr et al. [1998], Kosoresow and Hofmeyr [1997], Jagadish et al. [1999], Cabrera et al. [2001], Gonzalez and Dasgupta [2003], Dasgupta et al. [2000, 2002]. Ghosh et al. [1999a, 1998, 1999b], Debar et al. [1998], Eskin et al. [2001], Marceau [2000], Endler [1998], Lane et al. [1999, 1997b, 1997a]
Mixture of Models	Section 7.1.3	Eskin [2000]
Neural Networks	Section 4.1	Ghosh et al. [1998]
Support Vector Machines	Section 4.3	Hu et al. [2003], Heller et al. [2003]
Rule-Based Systems	Section 4.4	Lee et al. [1997, 1998, 2000]

different sequences. The length of the sequence for each program varies. Figure 5 illustrates a sample set of operating system call sequences. A key characteristic of the data in this domain is that the data can be typically profiled at different levels, such as program level or user level. Anomaly detection techniques applied for host-based intrusion detection are required to handle the sequential nature of data. Moreover, point anomaly detection techniques are not applicable in this domain. The techniques have to either model the sequence data or compute similarity between sequences. A survey of different techniques used for this problem is presented by Snyder [2001]. Comparative evaluations of anomaly detection for host-based intrusion detection are presented in Forrest et al. [1999] and Dasgupta and Nino [2000]. Some anomaly detection techniques used in this domain are shown in Table II.

3.1.2. Network Intrusion Detection Systems. These systems deal with detecting intrusions in network data. The intrusions typically occur as anomalous patterns (point anomalies) though certain techniques model the data in a sequential fashion and detect anomalous subsequences (collective anomalies) [Gwadera et al. 2005b, 2004]. The primary reason for these anomalies is due to the attacks launched by outside hackers who want to gain unauthorized access to the network for information theft or to disrupt the network. A typical setting is a large network of computers connected to the rest of the world via the Internet.

The data available for intrusion detection systems can be at different levels of granularity, for example, packet level traces, CISCO net-flows data, and so forth. The data has a temporal aspect associated with it but most of the techniques typically do not explicitly handle the sequential aspect. The data is high dimensional typically with a mix of categorical as well as continuous attributes.

A challenge faced by anomaly detection techniques in this domain is that the nature of anomalies keeps changing over time as the intruders adapt their network attacks to evade the existing intrusion detection solutions.

Some anomaly detection techniques used in this domain are shown in Table III.

3.2. Fraud Detection

Fraud detection refers to detection of criminal activities occurring in commercial organizations such as banks, credit card companies, insurance agencies, cell phone companies, stock market, and so on. The malicious users might be the actual customers of the organization or might be posing as customers (also known as *identity theft*). The fraud occurs when these users consume the resources provided by the organization in

Table III. Examples of Anomaly Detection Techniques Used for Network Intrusion Detection

Technique Used	Section	References
Statistical Profiling using Histograms	Section 7.2.1	NIDES Anderson et al. 1994, 1995; Javitz and Valdes 1991], EMERALD [Porras and Neumann 1997], Yamaniishi et al. [2001, 2004], Ho et al. [1999], Kruegel et al. [2002, 2003], Mahoney et al. [2003], Mahoney and Chan [2002, 2003]; Sargor [1998],
Parametric Statistical Modeling	Section 7.1	Gwadera et al. [2005b, 2004], Ye and Chen [2001], Tandon and Chan [2007]
Nonparametric Statistical Modeling	Section 7.2.2	Chow and Yeung [2002]
Bayesian Networks	Section 4.2	Siaterlis and Maglaris [2004], Sebyala et al. [2002], Valdes and Skinner [2000], Bronstein et al. [2001],
Neural Networks	Section 4.1	HIDE [Zhang et al. 2001], NSOM [Labib and Vemuri 2002], Smith et al. [2002], Hawkins et al. [2002], Kruegel et al. [2003], Manikopoulos and Papavassiliou [2002], Ramadas et al. [2003],
Support Vector Machines	Section 4.3	Eskin et al. [2002]
Rule-Based Systems	Section 4.4	ADAM [Barbara et al. 2001a, 2003, 2001b], Fan et al. [2001], Helmer et al. [1998], Qin and Hwang [2004], Salvador and Chan [2003], Otey et al. [2003]
Clustering Based	Section 6	ADMIT [Sequeira and Zaki 2002], Eskin et al. [2002], Wu and Zhang [2003], Otey et al. [2003]
Nearest Neighbor-Based	Section 5	MINDS [Ertoz et al. 2004; Chandola et al. 2006]
Spectral	Section 9	Shyu et al. [2003], Lakhina et al. [2005], Thottan and Ji [2003], Sun et al. [2007]
Information Theoretic	Section 8	Lee and Xiang [2001], Noble and Cook [2003]

an unauthorized way. The organizations are interested in immediate detection of such frauds to prevent economic losses.

Fawcett and Provost [1999] introduce the term *activity monitoring* as a general approach to fraud detection in these domains. The typical approach of anomaly detection techniques is to maintain a usage profile for each customer and monitor the profiles to detect any deviations. Some of the specific applications of fraud detection are discussed below.

3.2.1. Credit Card Fraud Detection. In this domain, anomaly detection techniques are applied to detect fraudulent credit card applications or fraudulent credit card usage (associated with credit card thefts). Detecting fraudulent credit card applications is similar to detecting insurance fraud [Ghosh and Reilly 1994].

The data is typically comprised of records defined over several dimensions such as user ID, amount spent, time between consecutive card usage, and so forth. The frauds are typically reflected in transactional records (point anomalies) and correspond to high payments, purchase of items never purchased by the user before, high rate of purchase, and so forth. The credit companies have complete data available and also have labeled records. Moreover, the data falls into distinct profiles based on the credit card user. Hence profiling and clustering based techniques are typically used in this domain.

The challenge associated with detecting unauthorized credit card usage is that it requires online detection of fraud as soon as the fraudulent transaction takes place.

Anomaly detection techniques have been applied in two different ways to address this problem. The first one is known as *by-owner* in which each credit card user is profiled based on his/her credit card usage history. Any new transaction is compared to the user's profile and flagged as an anomaly if it does not match the profile. This approach is typically expensive since it requires querying a central data repository, every time a user makes a transaction. Another approach known as *by-operation* detects anomalies

Table IV. Examples of Anomaly Detection Techniques Used for Credit Card Fraud Detection

Technique Used	Section	References
Neural Networks	Section 4.1	CARDWATCH [Aleskerov et al. 1997], Ghosh and Reilly [1994], Brause et al. [1999], Dorronsoro et al. [1997]
Rule-Based Systems	Section 4.4	Brause et al. [1999]
Clustering	Section 6	Bolton and Hand [1999]

Table V. Examples of Anomaly Detection Techniques Used for Mobile Phone Fraud Detection

Technique Used	Section	References
Statistical Profiling using Histograms	Section 7.2.1	Fawcett and Provost [1999], Cox et al. [1997]
Parametric Statistical Modeling	Section 7.1	Agarwal [2005], Scott [2001]
Neural Networks	Section 4.1	Barson et al. [1996], Taniguchi et al. [1998]
Rule-based Systems	Section 4.4	Phua et al. [2004], Taniguchi et al. [1998]

from among transactions taking place at a specific geographic location. Both *by-user* and *by-operation* techniques detect contextual anomalies. In the first case the context is a user, while in the second case the context is the geographic location.

Some anomaly detection techniques used in this domain are listed in Table IV.

3.2.2. Mobile Phone Fraud Detection. Mobile/cellular fraud detection is a typical activity monitoring problem. The task is to scan a large set of accounts, examining the calling behavior of each, and to issue an alarm when an account appears to have been misused.

Calling activity may be represented in various ways, but is usually described with call records. Each call record is a vector of features, both continuous (e.g., CALL-DURATION) and discrete (e.g., CALLING-CITY). However, there is no inherent primitive representation in this domain. Calls can be aggregated by time, for example into call-hours or call-days or user or area, depending on the granularity desired. The anomalies correspond to a high volume of calls or calls made to unlikely destinations.

Some techniques applied to cell phone fraud detection are listed in Table V.

3.2.3. Insurance Claim Fraud Detection. An important problem in the property-casualty insurance industry is claims fraud, for example, automobile insurance fraud. Individuals and conspiratorial rings of claimants and providers manipulate the claim processing system for unauthorized and illegal claims. Detection of such fraud has been very important for the associated companies to avoid financial losses.

The available data in this domain are the documents submitted by the claimants. The techniques extract different features, both categorical as well as continuous, from these documents. Typically, claim adjusters and investigators assess these claims for frauds. These manually investigated cases are used as labeled instances by supervised and semisupervised techniques for insurance fraud detection.

Insurance claim fraud detection is quite often handled as a generic activity monitoring problem [Fawcett and Provost 1999]. Neural network-based techniques have also been applied to identify anomalous insurance claims [He et al. 2003; Brockett et al. 1998].

3.2.4. Insider Trading Detection. Another recent application of anomaly detection techniques has been in early detection of *Insider Trading*. Insider trading is a phenomenon found in stock markets, where people make illegal profits by acting on (or leaking) inside information before the information is made public. The inside information can be of different forms [Donoho 2004]. It could refer to the knowledge of a

Table VI. Examples of Different Anomaly Detection Techniques Used for Insider Trading Detection

Technique Used	Section	References
Statistical Profiling using Histograms	Section 7.2.1	Donoho [2004], Aggarwal [2005]
Information Theoretic	Section 8	Arning et al. [1996]

Table VII. Examples of Different Anomaly Detection Techniques Used in Medical and Public Health Domain

Technique Used	Section	References
Parametric Statistical Modeling	Section 7.1	Horn et al. [2001], Laurikkala et al. [2000], Solberg and Lahti [2005], Roberts [2002], Suzuki et al. [2003]
Neural Networks	Section 4.1	Campbell and Bennett [2001]
Bayesian Networks	Section 4.2	Wong et al. [2003]
Rule-Based Systems	Section 4.4	Aggarwal [2005]
Nearest Neighbor based Techniques	Section 5	Lin et al. [2005]

pending merger/acquisition, a terrorist attack affecting a particular industry, pending legislation affecting a particular industry, or any information that would affect the stock prices in a particular industry. Insider trading can be detected by identifying anomalous trading activities in the market.

The available data is from several heterogeneous sources such as option trading data, stock trading data, news. The data has temporal associations since the data is collected continuously. The temporal and streaming nature has also been exploited in certain techniques [Aggarwal 2005].

Anomaly detection techniques in this domain are required to detect fraud in an online manner and as early as possible, to prevent people/organizations from making illegal profits.

Some anomaly detection techniques used in this domain are listed in Table VI.

3.3. Medical and Public Health Anomaly Detection

Anomaly detection in the medical and public health domains typically works with patient records. The data can have anomalies due to several reasons, such as abnormal patient condition, instrumentation errors, or recording errors. Several techniques have also focussed on detecting disease outbreaks in a specific area [Wong et al. 2003]. Thus the anomaly detection is a very critical problem in this domain and requires a high degree of accuracy.

The data typically consists of records that may have several different types of features, such as patient age, blood group, and weight. The data might also have a temporal as well as spatial aspect to it. Most of the current anomaly detection techniques in this domain aim at detecting anomalous records (point anomalies). Typically the labeled data belongs to the healthy patients, hence most of the techniques adopt a semisupervised approach. Another form of data handled by anomaly detection techniques in this domain is time-series data, such as *Electrocardiograms* (ECG) (Figure 4) and *Electroencephalograms* (EEG). Collective anomaly detection techniques have been applied to detect anomalies in such data [Lin et al. 2005].

The most challenging aspect of the anomaly detection problem in this domain is that the cost of classifying an anomaly as normal can be very high.

Some anomaly detection techniques used in this domain are listed in Table VII.

3.4. Industrial Damage Detection

Industrial units suffer damage due to continuous usage and normal wear and tear. Such damage needs to be detected early to prevent further escalation and losses. The

Table VIII. Examples of Anomaly Detection Techniques Used for Fault Detection in Mechanical Units

Technique Used	Section	References
Parametric Statistical Modeling	Section 7.1	Guttormsson et al. [1999], Keogh et al. [1997, 2002, 2006]
Non-parametric Statistical Modeling	Section 7.2.2	Desforges et al. [1998]
Neural Networks	Section 4.1	Bishop [1994], Campbell and Bennett [2001], Diaz and Hollmen [2002], Harris [1993], Jakubek and Strasser [2002], King et al. [2002], Li et al. [2002], Petsche et al. [1996], Streifel et al. [1996], Whitehead and Hoyt [1993]
Spectral	Section 9	Parra et al. [1996], Fujimaki et al. [2005]
Rule Based Systems	Section 4.4	Yairi et al. [2001]

Table IX. Examples of Anomaly Detection Techniques Used for Structural Damage Detection

Technique Used	Section	References
Statistical Profiling using Histograms	Section 7.2.1	Manson [2002], Manson et al. [2001, 2000]
Parametric Statistical Modeling	Section 7.1	Ruotolo and Surace [1997]
Mixture of Models	Section 7.1.3	Hickinbotham et al. [2000a, 2000b], Hollier and Austin [2002]
Neural Networks	Section 4.1	Brotherton et al. [1998, 2001], Nairac et al. [1999, 1997], Surace et al. [1998, 1997], Sohn et al. [2001], Worden [1997]

data in this domain is usually referred to as sensor data because it is recorded using different sensors and collected for analysis. Anomaly detection techniques have been extensively applied in this domain to detect such damage. Industrial damage detection can be further classified into two domains, one that deals with defects in mechanical components such as motors, engines, and so on, and the other that deals with defects in physical structures. The former domain is also referred to as *system health management*.

3.4.1. Fault Detection in Mechanical Units. The anomaly detection techniques in this domain monitor the performance of industrial components such as motors, turbines, oil flow in pipelines or other mechanical components, and detect defects that might occur due to wear and tear or other unforeseen circumstances.

The data in this domain typically has a temporal aspect and time-series analysis is also used in some techniques [Keogh et al. 2002, 2006; Basu and Meckesheimer 2007]. The anomalies occur mostly because of an observation in a specific context (contextual anomalies) or as an anomalous sequence of observations (collective anomalies).

Typically, normal data pertaining to components without defects is readily available and hence semisupervised techniques are applicable. Anomalies are required to be detected in an online fashion as preventive measures are required to be taken as soon as an anomaly occurs.

Some anomaly detection techniques used in this domain are listed in Table VIII.

3.4.2. Structural Defect Detection. Structural defect and damage detection techniques detect structural anomalies in structures, for example, cracks in beams or strains in airframes.

The data collected in this domain has a temporal aspect. The anomaly detection techniques are similar to novelty detection or change point detection techniques since they try to detect change in the data collected from a structure. The normal data and hence the models learned are typically static over time. The data might have spatial correlations.

Some anomaly detection techniques used in this domain are listed in Table IX.

Table X. Examples of Anomaly Detection Techniques Used in Image Processing Domain

Technique Used	Section	References
Mixture of Models	Section 7.1.3	Byers and Raftery [1998], Spence et al. [2001], Tarassenko [1995]
Regression	Section 7.1.2	Chen et al. [2005], Torr and Murray [1993]
Bayesian Networks	Section 4.2	Diehl and Hampshire [2002]
Support Vector Machines	Section 4.3	Davy and Godsill [2002], Song et al. [2002]
Neural Networks	Section 4.1	Augusteijn and Folkert [2002], Cun et al. [1990], Hazel [2000], Moya et al. [1993], Singh and Markou [2004]
Clustering	Section 6	Scarath et al. [1995]
Nearest Neighbor-Based Techniques	Section 5	Pokrajac et al. [2007], Byers and Raftery [1998]

Table XI. Examples of Anomaly Detection Techniques Used for Anomalous Topic Detection in Text Data

Technique Used	Section	References
Mixture of Models	Section 7.1.3	Baker et al. [1999]
Statistical Profiling using Histograms	Section 7.2.1	Fawcett and Provost [1999]
Support Vector Machines	Section 4.3	Manevitz and Yousef [2002]
Neural Networks	Section 4.1	Manevitz and Yousef [2000]
Clustering Based	Section 6	Allan et al. [1998], Srivastava and Zane-Ulman [2005], Srivastava [2006]

3.5. Image Processing

Anomaly detection techniques dealing with images are either interested in any changes in an image over time (motion detection) or in regions that appear abnormal on the static image. This domain includes satellite imagery [Augusteijn and Folkert 2002; Byers and Raftery 1998; Moya et al. 1993; Torr and Murray 1993; Theiler and Cai 2003], digit recognition [Cun et al. 1990], spectroscopy [Chen et al. 2005; Davy and Godsill 2002; Hazel 2000; Scarath et al. 1995], mammographic image analysis [Spence et al. 2001; Tarassenko 1995], and video surveillance [Diehl and Hampshire 2002; Singh and Markou 2004; Pokrajac et al. 2007]. The anomalies are caused by motion, or insertion of a foreign object, or instrumentation errors. The data has spatial as well as temporal characteristics. Each data point has a few continuous attributes such as color, lightness, texture, and so on. The interesting anomalies are either anomalous points or regions in the images (point and contextual anomalies).

One of the key challenges in this domain is the large size of the input. When dealing with video data, online anomaly detection techniques are required.

Some anomaly detection techniques used in this domain are listed in Table X.

3.6. Anomaly Detection in Text Data

Anomaly detection techniques in this domain primarily detect novel topics or events or news stories in a collection of documents or news articles. The anomalies are caused due to a new interesting event or an anomalous topic.

The data in this domain is typically high dimensional and very sparse. The data also has a temporal aspect since the documents are collected over time.

A challenge for anomaly detection techniques in this domain is to handle the large variations in documents belonging to one category or topic.

Some anomaly detection techniques used in this domain are listed in Table XI.

3.7. Sensor Networks

Sensor networks have lately become an important topic of research; more from the data analysis perspective, since the sensor data collected from various wireless sensors has

Table XII. Examples of Anomaly Detection Techniques Used for Anomaly Detection in Sensor Networks

Technique Used	Section	References
Bayesian Networks	Section 4.2	Janakiram et al. [2006]
Rule-Based Systems	Section 4.4	Branch et al. [2006]
Parametric Statistical Modeling	Section 7.1	Phuong et al. [2006], Du et al. [2006]
Nearest Neighbor-Based Techniques	Section 5	Subramaniam et al. [2006], Zhang et al. [2007], Idé et al. [2007]
Spectral	Section 9	Chatzigiannakis et al. [2006]

several unique characteristics. Anomalies in data collected from a sensor network can either mean that one or more sensors are faulty, or they are detecting events (such as intrusions) that are interesting for analysts. Thus anomaly detection in sensor networks can capture *sensor fault* detection or *intrusion* detection or both.

A single sensor network might be comprised of sensors that collect different types of data, such as binary, discrete, continuous, audio, video, and so forth. The data is generated in a streaming mode. Often times the environment in which the various sensors are deployed, and the communication channel, induce noise and missing values in the collected data.

Anomaly detection in sensor networks poses a set of unique challenges. The anomaly detection techniques are required to operate online. Due to severe resource constraints, the anomaly detection techniques need to be lightweight. Another challenge is that data is collected in a distributed fashion, and hence a distributed data mining approach is required to analyze it [Chatzigiannakis et al. 2006]. Moreover, the presence of noise in the data collected from the sensor makes anomaly detection more challenging, since it has to now distinguish between interesting anomalies and unwanted noise/missing values.

Table XII lists some anomaly detection techniques used in this domain.

3.8. Other Domains

Anomaly detection has also been applied to several other domains such as speech recognition [Albrecht et al. 2000; Emamian et al. 2000], novelty detection in robot behavior [Crook and Hayes 2001; Crook et al. 2002; Marsland et al. 1999, 2000b, 2000a], traffic monitoring [Shekhar et al. 2001], click-through protection [Ihler et al. 2006], detecting faults in Web applications [Ide and Kashima 2004; Sun et al. 2005], detecting anomalies in biological data [Kadota et al. 2003; Sun et al. 2006; Gwadera et al. 2005a; MacDonald and Ghosh 2007; Tomlins et al. 2005; Tibshirani and Hastie 2007], detecting anomalies in census data [Lu et al. 2003], detecting associations among criminal activities [Lin and Brown 2003], detecting anomalies in *Customer Relationship Management* (CRM) data [He et al. 2004b], detecting anomalies in astronomical data [Dutta et al. 2007; Escalante 2005; Protopapas et al. 2006] and detecting ecosystem disturbances [Blender et al. 1997; Kou et al. 2006; Sun and Chawla 2004].

4. CLASSIFICATION BASED ANOMALY DETECTION TECHNIQUES

Classification [Tan et al. 2005; Duda et al. 2000] is used to learn a model (*classifier*) from a set of labeled data instances (*training*) and then, classify a test instance into one of the classes using the learned model (*testing*). Classification-based anomaly detection techniques operate in a similar two-phase fashion. The training phase learns a classifier using the available labeled training data. The testing phase classifies a test instance as normal or anomalous, using the classifier.

Classification based anomaly detection techniques operate under the following general assumption:

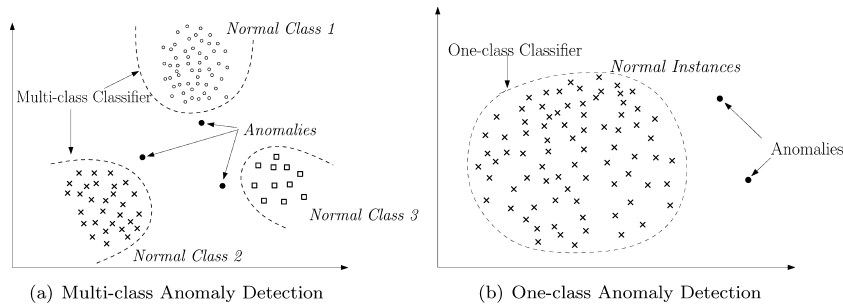


Fig. 6. Using classification for anomaly detection.

Assumption. A classifier that can distinguish between normal and anomalous classes can be learned in the given feature space.

Based on the labels available for the training phase, classification-based anomaly detection techniques can be grouped into two broad categories: *multi-class* and *one-class* anomaly detection techniques.

Multi-class classification based anomaly detection techniques assume that the training data contains labeled instances belonging to multiple normal classes [Stefano et al. 2000; Barbara et al. 2001b]. Such anomaly detection techniques teach a classifier to distinguish between each normal class and the rest of the classes. See Figure 6(a) for illustration. A test instance is considered anomalous if it is not classified as normal by any of the classifiers. Some techniques in this subcategory associate a confidence score with the prediction made by the classifier. If none of the classifiers are confident in classifying the test instance as normal, the instance is declared to be anomalous.

One-class classification based anomaly detection techniques assume that all training instances have only one class label. Such techniques learn a discriminative boundary around the normal instances using a *one-class classification algorithm*, for example, one-class SVMs [Schölkopf et al. 2001], one-class Kernel Fisher Discriminants [Roth 2004, 2006], as shown in Figure 6(b). Any test instance that does not fall within the learned boundary is declared as anomalous.

In the following subsections, we discuss a variety of anomaly detection techniques that use different classification algorithms to build classifiers:

4.1. Neural Networks-Based

Neural networks have been applied to anomaly detection in multi-class as well as one-class settings.

A basic multi-class anomaly detection technique using neural networks operates in two steps. First, a neural network is trained on the normal training data to learn the different normal classes. Second, each test instance is provided as an input to the neural network. If the network accepts the test input, it is normal and if the network rejects a test input, it is an anomaly [Stefano et al. 2000; Odin and Addison 2000]. Several variants of the basic neural network technique have been proposed that use different types of neural networks, as summarized in Table XIII.

Replicator Neural Networks have been used for one-class anomaly detection [Hawkins et al. 2002; Williams et al. 2002]. A multi-layer feed forward neural network is constructed that has the same number of input and output neurons (corresponding to the features in the data). The training involves compressing data into three hidden layers. The testing phase involves reconstructing each data instance, x_i , using the learned

Table XIII. Some Examples of Classification-Based Anomaly Detection Techniques Using Neural Networks

Neural Network Used	References
Multi layered Perceptrons	[Augusteijn and Folkert 2002; Cun et al. 1990; Sykacek 1997; Ghosh et al. 1999a, 1998; Barson et al. 1996; He et al. 1997; Nairac et al. 1997; Hickinbotham and Austin 2000b; Vasconcelos et al. 1995, 1994]
Neural Trees	[Martinez 1998]
Autoassociative Networks	[Aeyels 1991; Byungho and Sungzoon 1999; Japkowicz et al. 1995; Hawkins et al. 2002; Ko and Jacyna 2000; Manevitz and Yousef 2000; Petsche et al. 1996; Sohn et al. 2001; Song et al. 2001; Streifel et al. 1996; Thompson et al. 2002; Worden 1997; Williams et al. 2002; Diaz and Hollmen 2002]
Adaptive Resonance Theory Based	[Moya et al. 1993; Dasgupta and Nino 2000; Caudell and Newman 1993]
Radial Basis Function-Based	[Albrecht et al. 2000; Bishop 1994; Brotherton et al. 1998; Brotherton and Johnson 2001; Li et al. 2002; Nairac et al. 1999, 1997; Ghosh and Reilly 1994; Jakubek and Strasser 2002]
Hopfield Networks	[Jagota 1991; Crook and Hayes 2001; Crook et al. 2002; Addison et al. 1999; Murray 2001]
Oscillatory Networks	Ho and Rouat 1997, 1998; Kojima and Ito 1999; Borisjuk et al. 2000; Martinelli and Perfetti 1994]

network to obtain the reconstructed output, o_i . The reconstruction error, δ_i , for the test instance x_i is then computed as:

$$\delta_i = \frac{1}{n} \sum_{j=1}^n (x_{ij} - o_{ij})^2,$$

where n is the number of features over which the data is defined. The reconstruction error δ_i is directly used as an anomaly score for the test instance.

4.2. Bayesian Networks-Based

Bayesian networks have been used for anomaly detection in the multi-class setting. A basic technique for a univariate categorical data set using a naïve Bayesian network estimates the posterior probability of observing a class label from a set of normal class labels and the anomaly class label, given a test data instance. The class label with largest posterior is chosen as the predicted class for the given test instance. The likelihood of observing the test instance given a class and the prior on the class probabilities, is estimated from the training data set. The zero probabilities, especially for the anomaly class, are smoothed using *Laplace Smoothing*.

The basic technique can be generalized to multivariate categorical data sets by aggregating the per-attribute posterior probabilities for each test instance and using the aggregated value to assign a class label to the test instance.

Several variants of the basic technique have been proposed for network intrusion detection [Barbara et al. 2001b; Sebyala et al. 2002; Valdes and Skinner 2000; Mingming 2000; Bronstein et al. 2001], for novelty detection in video surveillance [Diehl and Hampshire 2002], for anomaly detection in text data [Baker et al. 1999], and for disease outbreak detection [Wong et al. 2002, 2003].

This basic technique assumes independence between the different attributes. Several variations of the basic technique have been proposed that capture the conditional dependencies between the different attributes using more complex Bayesian networks [Siaterlis and Maglaris 2004; Janakiram et al. 2006; Das and Schneider 2007].

4.3. Support Vector Machines-Based

Support Vector Machines (SVMs) [Vapnik 1995] have been applied to anomaly detection in the one-class setting. Such techniques use one class learning techniques for SVM [Ratsch et al. 2002] and learn a region that contains the training data instances (a boundary). Kernels, such as *radial basis function (RBF) kernel*, can be used to learn complex regions. For each test instance, the basic technique determines if the test instance falls within the learned region. If a test instance falls within the learned region, it is declared as normal, else it is declared as anomalous.

Variants of the basic technique have been proposed for anomaly detection in audio signal data [Davy and Godsill 2002], novelty detection in power generation plants [King et al. 2002] and system call intrusion detection [Eskin et al. 2002; Heller et al. 2003; Lazarevic et al. 2003]. The basic technique has also been extended to detect anomalies in temporal sequences [Ma and Perkins 2003a, 2003b].

A variant of the basic technique [Tax and Duin 1999a, 1999b; Tax 2001] finds the smallest hypersphere in the kernel space that contains all training instances, and then determines on which side of that hypersphere a test instance lies. If a test instance lies outside the hypersphere, it is declared to be anomalous.

Song et al. [2002] use *Robust Support Vector Machines (RSVM)*, which are robust to the presence of anomalies in the training data. RSVM have been applied to system call intrusion detection [Hu et al. 2003].

4.4. Rule-Based

Rule-based anomaly detection techniques learn rules that capture the normal behavior of a system. A test instance that is not covered by any such rule is considered as an anomaly. Rule-based techniques have been applied in multi-class as well as one-class settings.

A basic multi-class rule-based technique consists of two steps. The first step is to learn rules from the training data using a rule learning algorithm, such as RIPPER, Decision Trees, and so on. Each rule has an associated confidence value that is proportional to ratio between the number of training instances correctly classified by the rule and the total number of training instances covered by the rule. The second step is to find, for each test instance, the rule that best captures the test instance. The inverse of the confidence associated with the best rule is the anomaly score of the test instance. Several minor variants of the basic rule-based technique have been proposed [Fan et al. 2001; Helmer et al. 1998; Lee et al. 1997; Salvador and Chan 2003; Teng et al. 1990].

Association rule mining [Agrawal and Srikant 1995] has been used for one-class anomaly detection by generating rules from the data in an unsupervised fashion. Association rules are generated from a categorical data set. To ensure that the rules correspond to strong patterns, a support threshold is used to prune out rules with low support [Tan et al. 2005]. Association rule mining-based techniques have been used for network intrusion detection [Mahoney and Chan 2002, 2003; Mahoney et al. 2003; Tandon and Chan 2007; Barbara et al. 2001a; Otey et al. 2003], system call intrusion detection [Lee et al. 2000; Lee and Stolfo 1998; Qin and Hwang 2004], credit card fraud detection [Brause et al. 1999], and fraud detection in spacecraft housekeeping data [Yairi et al. 2001]. Frequent itemsets are generated in the intermediate step of association rule mining algorithms. He et al. [2004a] propose an anomaly detection algorithm for categorical data sets in which the anomaly score of a test instance is equal to the number of frequent itemsets in which it occurs.

Computational Complexity. The computational complexity of classification-based techniques depends on the classification algorithm being used. For a discussion on

the complexity of training classifiers, see Kearns [1990]. Generally, training decision trees tends to be faster, while techniques that involve quadratic optimization, such as SVMs, are more expensive; though linear time SVMs [Joachims 2006] have been proposed that have linear training time. The testing phase of classification techniques is usually very fast since the testing phase uses a learned model for classification.

Advantages and Disadvantages of Classification-Based Techniques. The advantages of classification-based techniques are as follows:

- (1) Classification-based techniques, especially the multi-class techniques, can make use of powerful algorithms that can distinguish between instances belonging to different classes.
- (2) The testing phase of classification-based techniques is fast, since each test instance needs to be compared against the precomputed model.

The disadvantages of classification-based techniques are as follows:

- (1) Multi-class classification-based techniques rely on the availability of accurate labels for various normal classes, which is often not possible.
- (2) Classification-based techniques assign a label to each test instance, which can also become a disadvantage when a meaningful anomaly score is desired for the test instances. Some classification techniques that obtain a probabilistic prediction score from the output of a classifier, can be used to address this issue [Platt 2000].

5. NEAREST NEIGHBOR-BASED ANOMALY DETECTION TECHNIQUES

The concept of nearest neighbor analysis has been used in several anomaly detection techniques. Such techniques are based on the following key assumption:

Assumption. Normal data instances occur in dense neighborhoods, while anomalies occur far from their closest neighbors.

Nearest neighbor-based anomaly detection techniques require a distance or similarity measure defined between two data instances. Distance (or similarity) between two data instances can be computed in different ways. For continuous attributes, Euclidean distance is a popular choice, but other measures can be used [Tan et al. 2005, Chapter 2]. For categorical attributes, a simple matching coefficient is often used but more complex distance measures can also be used [Boriah et al. 2008; Chandola et al. 2008]. For multivariate data instances, distance or similarity is usually computed for each attribute and then combined [Tan et al. 2005, Chapter 2].

Most of the techniques that will be discussed in this section, as well as the clustering-based techniques (Section 6) do not require the distance measure to be strictly metric. The measures are typically required to be *positive-definite* and *symmetric*, but they are not required to satisfy the *triangle inequality*.

Nearest neighbor-based anomaly detection techniques can be broadly grouped into two categories:

- (1) techniques that use the distance of a data instance to its k^{th} nearest neighbor as the anomaly score;
- (2) techniques that compute the relative density of each data instance to compute its anomaly score.

Additionally there are some techniques that use the distance between data instances in a different manner to detect anomalies, and will be briefly discussed later.

5.1. Using Distance to k^{th} Nearest Neighbor

A basic nearest neighbor anomaly detection technique is based on the following definition: *The anomaly score of a data instance is defined as its distance to its k^{th} nearest neighbor in a given data set.* This basic technique has been applied to detect land mines from satellite ground images [Byers and Raftery 1998] and to detect shorted turns (anomalies) in the DC field windings of large synchronous turbine-generators [Guttormsson et al. 1999]. In the latter paper the authors use $k = 1$. Usually, a threshold is then be applied on the anomaly score to determine if a test instance is anomalous or not. Ramaswamy et al. [2000], on the other hand, select n instances with the largest anomaly scores as the anomalies.

The basic technique has been extended by researchers in three different ways. The first set of variants modifies the definition to obtain the anomaly score of a data instance. The second set of variants uses different distance/similarity measures to handle different data types. The third set of variants focuses on improving the efficiency of the basic technique (the complexity of the basic technique is $\mathcal{O}(N^2)$, where N is the data size) in different ways.

Eskin et al. [2002], Angiulli and Pizzuti [2002] and Zhang and Wang [2006] compute the anomaly score of a data instance as the sum of its distances from its k nearest neighbors. A similar technique, called Peer Group Analysis has been applied to detect credit card frauds by Bolton and Hand [1999].

A different way to compute the anomaly score of a data instance is to count the number of nearest neighbors (n) that are not more than d distance apart from the given data instance [Knorr and Ng 1997, 1998, 1999; Knorr et al. 2000]. This method can also be viewed as estimating the global density for each data instance, since it involves counting the number of neighbors in a hypersphere of radius d . For example, in a 2-D data set, the density of a data instance $= \frac{n}{\pi d^2}$. The inverse of the density is the anomaly score for the data instance. Instead of computing the actual density, several techniques fix the radius d and use $\frac{1}{n}$ as the anomaly score, while several techniques fix n and use $\frac{1}{d}$ as the anomaly score.

While most techniques discussed so far in this category have been proposed to handle continuous attributes, several variants have been proposed to handle other data types. A hypergraph-based technique, called HOT, is proposed by Wei et al. [2003] in which the authors model the categorical values using a hypergraph, and measure distance between two data instances by analyzing the connectivity of the graph. A distance measure for data containing a mix of categorical and continuous attributes has been proposed for anomaly detection [Otey et al. 2006]. The authors define links between two instances by adding distance for categorical and continuous attributes separately. For categorical attributes, the number of attributes for which the two instances have the same values defines the distance between them. For continuous attributes, a covariance matrix is maintained to capture the dependencies between the continuous values. Palshikar [2005] adapts the technique proposed in Knorr and Ng [1999] to continuous sequences. Kou et al. [2006] extend the technique proposed in Ramaswamy et al. [2000] to spatial data.

Several variants of the basic technique have been proposed to improve the efficiency. Some techniques prune the search space by either ignoring instances that cannot be anomalous or by focussing on instances that are most likely to be anomalous. Bay and Schwabacher [2003] show that for sufficiently randomized data, a simple pruning step could result in the average complexity of the nearest neighbor search to be nearly linear. After calculating the nearest neighbors for a data instance, the algorithm sets the anomaly threshold for any data instance to the score of the weakest anomaly found

so far. Using this pruning procedure, the technique discards instances that are close, and hence not interesting.

Ramaswamy et al. [2000] propose a partition-based technique, which first clusters the instances and computes lower and upper bounds on the distance of an instance from its k^{th} nearest neighbor for instances in each partition. This information is then used to identify the partitions that cannot possibly contain the top k anomalies; such partitions are pruned. Anomalies are then computed from the remaining instances (belonging to unpruned partitions) in a final phase. Similar cluster-based pruning has been proposed by Eskin et al. [2002], McCallum et al. [2000], Ghoting et al. [2006], and Tao et al. [2006].

Wu and Jermaine [2006] use sampling to improve the efficiency of the nearest neighbor-based technique. The authors compute the nearest neighbor of every instance within a smaller sample from the data set. Thus the complexity of the proposed technique is reduced to $\mathcal{O}(MN)$, where M is the sample size chosen.

To prune the search space for nearest neighbors, several techniques partition the attribute space into a hypergrid consisting of hypercubes of fixed sizes. The intuition behind such techniques is that if a hypercube contains many instances, such instances are likely to be normal. Moreover, if for a given instance, the hypercube that contains the instance, and its adjoining hypercubes, contain very few instances, the given instance is likely to be anomalous. Techniques based on this intuition have been proposed by Knorr and Ng [1998]. Angiulli and Pizzuti [2002] extend by linearizing the search space through the Hilbert space filling curve. The d -dimensional data set is fitted in a hypercube $D = [0, 1]^d$. This hypercube is then mapped to the interval $I = [0, 1]$ using the *Hilbert Space Filling Curve* and the k -nearest neighbors of a data instance are obtained by examining its successors and predecessors in I .

5.2. Using Relative Density

Density-based anomaly detection techniques estimate the density of the neighborhood of each data instance. An instance that lies in a neighborhood with low density is declared to be anomalous while an instance that lies in a dense neighborhood is declared to be normal.

For a given data instance, the distance to its k^{th} nearest neighbor is equivalent to the radius of a hypersphere, centered at the given data instance, which contains k other instances. Thus the distance to the k^{th} nearest neighbor for a given data instance can be viewed as an estimate of the inverse of the density of the instance in the data set and the basic nearest neighbor-based technique described in the previous subsection can be considered as a density-based anomaly detection technique.

Density-based techniques perform poorly if the data has regions of varying densities. For example, consider the two-dimensional data set shown in Figure 7. Due to the low density of the cluster C_1 , it is apparent that for every instance q inside the cluster C_1 , the distance between the instance q and its nearest neighbor is greater than the distance between the instance p_2 and the nearest neighbor from the cluster C_2 , and the instance p_2 will not be considered as anomaly. Hence, the basic technique will fail to distinguish between p_2 and instances in C_1 . However, the instance p_1 may be detected.

To handle the issue of varying densities in the data set, a set of techniques has been proposed to compute the density of instances relative to the density of their neighbors.

Breunig et al. [1999, 2000] assign an anomaly score to a given data instance, known as *Local Outlier Factor* (LOF). For any given data instance, the LOF score is equal to ratio of average local density of the k nearest neighbors of the instance and the local density of the data instance itself. To find the local density for a data instance, the authors first find the radius of the smallest hyper-sphere centered at the data instance, that contains its k nearest neighbors. The local density is then computed by dividing k

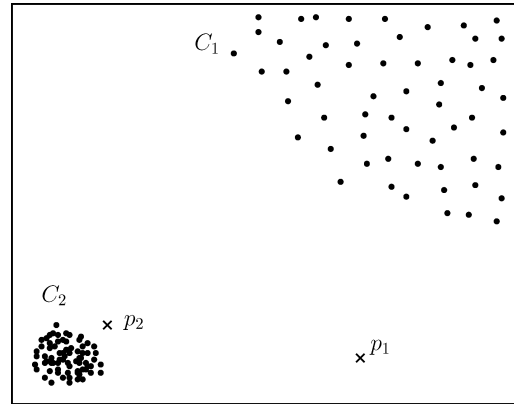


Fig. 7. Advantage of local density-based techniques over global density-based techniques.

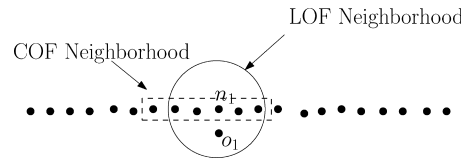


Fig. 8. Difference between the neighborhoods computed by LOF and COF.

by the volume of this hyper-sphere. For a normal instance lying in a dense region, its local density will be similar to that of its neighbors, while for an anomalous instance, its local density will be lower than that of its nearest neighbors. Hence the anomalous instance will get a higher LOF score.

In the example shown in Figure 7, LOF will be able to capture both anomalies, p_1 and p_2 , due to the fact that it considers the density around the data instances.

Several researchers have proposed variants of the LOF technique. Some of these variants estimate the local density of an instance in a different way. Some variants have adapted the original technique to more complex data types. Since the original LOF technique is $\mathcal{O}(N^2)$ (N is the data size), several techniques have been proposed that improve the efficiency of LOF.

Tang et al. [2002] discuss a variation of the LOF, which they call *Connectivity-based Outlier Factor* (COF). The difference between LOF and COF is the manner in which the k neighborhood for an instance is computed. In COF, the neighborhood for an instance is computed in an incremental mode. To start, the closest instance to the given instance is added to the neighborhood set. The next instance added to the neighborhood set is such that its distance to the existing neighborhood set is minimum among all remaining data instances. The distance between an instance and a set of instances is defined as the minimum distance between the given instance and any instance belonging to the given set. The neighborhood is grown in this manner until it reaches size k . Once the neighborhood is computed, the anomaly score (COF) is computed in the same manner as LOF. COF is able to capture regions such as straight lines, as shown in Figure 8.

A simpler version of LOF was proposed by Hautamaki et al. [2004], which calculates a quantity called *Outlier Detection using In-degree Number* (ODIN) for each data instance. For a given data instance, ODIN is equal to the number of k nearest neighbors of the data instance which have the given data instance in their k nearest neighbor list.

The inverse of ODIN is the anomaly score for the data instance. A similar technique was proposed by Brito et al. [1997].

Papadimitriou et al. [2002] propose a measure called *Multi-Granularity Deviation Factor* (MDEF), which is a variation of LOF. MDEF for a given data instance is equal to the standard deviation of the local densities of the nearest neighbors of the given data instance (including the data instance itself). The inverse of the standard deviation is the anomaly score for the data instance. The anomaly detection technique presented in the paper is called LOCI, which not only finds anomalous instances but also anomalous micro-clusters.

Several variants of LOF have been proposed to handle different data types. A variant of LOF is applied by Sun and Chawla [2004, 2006] for detecting spatial anomalies in climate data. Yu et al. [2006] use a similarity measure instead of distance to handle categorical attributes. A similar technique has been proposed to detect sequential anomalies in protein sequences by Sun et al. [2006]. This technique uses *Probabilistic Suffix Trees* (PST) to find the nearest neighbors for a given sequence. Pokrajac et al. [2007] extend LOF to work in an incremental fashion to detect anomalies in video sensor data.

Some variants of the LOF technique have been proposed to improve its efficiency. Jin et al. [2001] propose a variant in which only the top n anomalies are found instead of finding LOF score for every data instance. The technique includes finding micro-clusters in the data and then finding upper and lower bounds on LOF for each of the micro-clusters. Chiu and Chee Fu [2003] proposed three variants of LOF, which enhance its performance by making certain assumptions about the problem to prune all those clusters that definitely do not contain instances that will figure in the top n anomaly list. For the remaining clusters a detailed analysis is done to find the LOF score for each instance in these clusters.

Computational Complexity. A drawback of the basic nearest neighbor-based technique and of the LOF technique, is the $O(N^2)$ complexity required. Since these techniques involve finding nearest neighbors for each instance, efficient data structures such as k - d trees [Bentley 1975] and R-trees [Roussopoulos et al. 1995] can be used. But such techniques do not scale well as the number of attributes increases. Several techniques have directly optimized the anomaly detection technique under the assumption that only the top few anomalies are interesting. If an anomaly score is required for every test instance, such techniques are not applicable. Techniques that partition the attribute space into a hypergrid, are linear in data size but are exponential in the number of attributes, and hence are not well suited for large number of attributes. Sampling techniques try to address the $O(N^2)$ complexity issue by determining the nearest neighbors within a small sample of the data set. But sampling might result in incorrect anomaly scores if the size of the sample is very small.

Advantages and Disadvantages of Nearest Neighbor-Based Techniques. The advantages of nearest neighbor-based techniques are as follows:

- (1) A key advantage of nearest neighbor-based techniques is that they are unsupervised in nature and do not make any assumptions regarding the generative distribution for the data. Instead, they are purely data driven.
- (2) Semisupervised techniques perform better than unsupervised techniques in terms of missed anomalies, since the likelihood that an anomaly will form a close neighborhood in the training data set is very low.
- (3) Adapting nearest neighbor-based techniques to a different data type is straightforward, and primarily requires defining an appropriate distance measure for the given data.

The disadvantages of nearest neighbor-based techniques are as follows:

- (1) For unsupervised techniques, if the data has normal instances that do not have enough close neighbors, or if the data has anomalies that have enough close neighbors, the technique fails to label them correctly; resulting in missed anomalies.
- (2) For semisupervised techniques, if the normal instances in the test data do not have enough similar normal instances in the training data, the false positive rate for such techniques is high.
- (3) The computational complexity of the testing phase is also a significant challenge since it involves computing the distance of each test instance—with all instances belonging to either the test data itself, or to the training data—to compute the nearest neighbors.
- (4) Performance of a nearest neighbor-based technique greatly relies on a distance measure, defined between a pair of data instances, which can effectively distinguish between normal and anomalous instances. Defining distance measures between instances can be challenging when the data is complex, for example, graphs, sequences, and so on.

6. CLUSTERING-BASED ANOMALY DETECTION TECHNIQUES

Clustering [Jain and Dubes 1988; Tan et al. 2005] is used to group similar data instances into clusters. Clustering is primarily an unsupervised technique though semisupervised clustering [Basu et al. 2004] has also been explored lately. Even though clustering and anomaly detection appear to be fundamentally different from each other, several clustering-based anomaly detection techniques have been developed. Clustering-based anomaly detection techniques can be grouped into three categories.

The first category of clustering-based techniques relies on the following assumption:

Assumption. Normal data instances belong to a cluster in the data, while anomalies do not belong to any cluster.

Techniques based on this assumption apply a known clustering-based algorithm to the data set and declare any data instance that does not belong to any cluster as anomalous. Several clustering algorithms that do not force every data instance to belong to a cluster, such as DBSCAN [Ester et al. 1996], ROCK [Guha et al. 2000], and SNN clustering [Ertöz et al. 2003] can be used. The *FindOut* algorithm [Yu et al. 2002] is an extension of the *WaveCluster* algorithm [Sheikholeslami et al. 1998] in which the detected clusters are removed from the data and the residual instances are declared as anomalies.

A disadvantage of such techniques is that they are not optimized to find anomalies, since the main aim of the underlying clustering algorithm is to find clusters.

The second category of clustering-based techniques relies on the following assumption:

Assumption. Normal data instances lie close to their closest cluster centroid, while anomalies are far away from their closest cluster centroid.

Techniques based on this assumption consist of two steps. In the first step, the data is clustered using a clustering algorithm. In the second step, for each data instance, its distance to its closest cluster centroid is calculated as its anomaly score.

A number of anomaly detection techniques that follow this two step approach have been proposed using different clustering algorithms. Smith et al. [2002] studied *Self-Organizing Maps* (SOM), K-means Clustering, and *Expectation Maximization* (EM) to

cluster training data and then use the clusters to classify test data. In particular, SOM [Kohonen 1997] has been widely used to detect anomalies in a semisupervised mode in several applications such as intrusion detection [Labib and Vemuri 2002; Smith et al. 2002; Ramadas et al. 2003], fault detection [Harris 1993; Ypma and Duin 1998; Emamian et al. 2000], and fraud detection [Brockett et al. 1998]. Barbara et al. [2003] propose a technique that is robust to anomalies in the training data. The authors first separate normal instances from anomalies in the training data, using frequent item-set mining, and then use the clustering-based technique to detect anomalies. Several techniques have also been proposed to handle sequence data [Blender et al. 1997; Bejerano and Yona 2001; Vinueza and Grudic 2004; Budalakoti et al. 2006].

Techniques based on the second assumption can also operate in semisupervised mode, in which the training data is clustered and instances belonging to the test data are compared against the clusters to obtain an anomaly score for the test data instance [Marchette 1999; Wu and Zhang 2003; Vinueza and Grudic 2004; Allan et al. 1998]. If the training data has instances belonging to multiple classes, semisupervised clustering can be applied to improve the clusters. He et al. [2002] incorporate the knowledge of labels to improve on their unsupervised clustering-based anomaly detection technique [He et al. 2003] by calculating a measure called *semantic anomaly factor*, which is high if the class label of an object in a cluster is different from the majority of the class labels in that cluster.

Note that if the anomalies in the data form clusters by themselves, these techniques will not be able to detect such anomalies. To address this issue, a third category of clustering-based techniques has been proposed, which relies on the following assumption:

Assumption. Normal data instances belong to large and dense clusters, while anomalies either belong to small or sparse clusters.

Techniques based on this assumption declare instances belonging to clusters whose size and/or density is below a threshold, as anomalous.

Several variations of the third category of techniques have been proposed [Pires and Santos-Pereira 2005; Otey et al. 2003; Eskin et al. 2002; Mahoney et al. 2003; Jiang et al. 2001; He et al. 2003]. The technique proposed by He et al. [2003], called *FindCBLOF*, assigns an anomaly score known as *Cluster-Based Local Outlier Factor* (CBLOF) for each data instance. The CBLOF score captures the size of the cluster to which the data instance belongs, as well as the distance of the data instance to its cluster centroid.

Several clustering-based techniques have been proposed to improve the efficiency of these existing techniques. Fixed width clustering is a linear time ($\mathcal{O}(Nd)$) approximation algorithm used by various anomaly detection techniques [Eskin et al. 2002; Portnoy et al. 2001; Mahoney et al. 2003; He et al. 2003]. An instance is assigned to a cluster whose center is within a pre-specified distance to the given instance. If no such cluster exists then a new cluster with the instance as the center is created. Then they determine which clusters are anomalies based on their density and distance from other clusters. The width can either be a user-specified parameter [Eskin et al. 2002; Portnoy et al. 2001] or can be derived from the data [Mahoney et al. 2003]. Chaudhary et al. [2002] propose an anomaly detection technique using k - d trees that provide a partitioning of the data in linear time. They apply their technique to detect anomalies in astronomical data sets where computational efficiency is an important requirement. Another technique that addresses this issue is proposed by Sun et al. [2004]. The authors propose an indexing technique called *CD-trees* to efficiently partition data into clusters. The data instances that belong to sparse clusters are declared as anomalies.

6.1. Distinction between Clustering-Based and Nearest Neighbor-Based Techniques

Several clustering-based techniques require distance computation between a pair of instances. Thus, in that respect, they are similar to nearest neighbor-based techniques. The choice of the distance measure is critical to the performance of the technique; hence the discussion in the previous section regarding the distance measures holds for clustering-based techniques also. The key difference between the two techniques, however, is that clustering-based techniques evaluate each instance with respect to the cluster it belongs to, while nearest neighbor-based techniques analyze each instance with respect to its local neighborhood.

Computational Complexity. The computational complexity of training a clustering-based anomaly detection technique depends on the clustering algorithm used to generate clusters from the data. Thus such techniques can have quadratic complexity if the clustering technique requires computation of pairwise distances for all data instances, or linear when using heuristic-based techniques such as k -means [Hartigan and Wong 1979] or approximate clustering techniques [Eskin et al. 2002]. The test phase of clustering-based techniques is fast, since it involves comparing a test instance with a small number of clusters.

Advantages and Disadvantages of Clustering-Based Techniques. The advantages of clustering based techniques are as follows:

- (1) Clustering-based techniques can operate in an unsupervised mode.
- (2) Such techniques can often be adapted to other complex data types by simply plugging in a clustering algorithm that can handle the particular data type.
- (3) The testing phase for clustering-based techniques is fast since the number of clusters against which every test instance needs to be compared is a small constant.

The disadvantages of clustering-based techniques are as follows:

- (1) Performance of clustering-based techniques is highly dependent on the effectiveness of clustering algorithms in capturing the cluster structure of normal instances.
- (2) Many techniques detect anomalies as a byproduct of clustering, and hence are not optimized for anomaly detection.
- (3) Several clustering algorithms force every instance to be assigned to some cluster. This might result in anomalies getting assigned to a large cluster, thereby being considered as normal instances by techniques that operate under the assumption that anomalies do not belong to any cluster.
- (4) Several clustering-based techniques are effective only when the anomalies do not form significant clusters among themselves.
- (5) The computational complexity for clustering the data is often a bottleneck, especially if $\mathcal{O}(N^2d)$ clustering algorithms are used.

7. STATISTICAL ANOMALY DETECTION TECHNIQUES

The underlying principle of any statistical anomaly detection technique is: “An anomaly is an observation which is suspected of being partially or wholly irrelevant because it is not generated by the stochastic model assumed” [Anscombe and Guttman 1960]. Statistical anomaly detection techniques are based on the following key assumption:

Assumption. Normal data instances occur in high probability regions of a stochastic model, while anomalies occur in the low probability regions of the stochastic model.

Statistical techniques fit a statistical model (usually for normal behavior) to the given data and then apply a statistical inference test to determine if an unseen instance belongs to this model or not. Instances that have a low probability of being generated from the learned model, based on the applied test statistic, are declared as anomalies. Both parametric as well as nonparametric techniques have been applied to fit a statistical model. While parametric techniques assume the knowledge of the underlying distribution and estimate the parameters from the given data [Eskin 2000], nonparametric techniques do not generally assume knowledge of the underlying distribution [Desforges et al. 1998]. In the next two subsections we will discuss parametric and nonparametric anomaly detection techniques.

7.1. Parametric Techniques

As mentioned before, parametric techniques assume that the normal data is generated by a parametric distribution with parameters Θ and probability density function $f(\mathbf{x}, \Theta)$, where $\mathbf{x}$ is an observation. The anomaly score of a test instance (or observation) $\mathbf{x}$ is the inverse of the probability density function, $f(\mathbf{x}, \Theta)$. The parameters Θ are estimated from the given data.

Alternatively, a statistical hypothesis test (also referred to as *discordancy test* in statistical outlier detection literature [Barnett and Lewis 1994]) may be used. The null hypothesis (H_0) for such tests is that the data instance $\mathbf{x}$ has been generated using the estimated distribution (with parameters Θ). If the statistical test rejects H_0 , $\mathbf{x}$ is declared to be anomaly. A statistical hypothesis test is associated with a test statistic, which can be used to obtain a probabilistic anomaly score for the data instance $\mathbf{x}$.

Based on the type of distribution assumed, parametric techniques can be further categorized as follows:

7.1.1. Gaussian Model-Based. Such techniques assume that the data is generated from a Gaussian distribution. The parameters are estimated using *Maximum Likelihood Estimates* (MLE). The distance of a data instance to the estimated mean is the anomaly score for that instance. A threshold is applied to the anomaly scores to determine the anomalies. Different techniques in this category calculate the distance to the mean and the threshold in different ways.

A simple outlier detection technique, often used in process quality control domain [Shewhart 1931], is to declare all data instances that are more than 3σ distance away from the distribution mean μ , where σ is the standard deviation for the distribution. The $\mu \pm 3\sigma$ region contains 99.7% of the data instances.

More sophisticated statistical tests have also been used to detect anomalies, as discussed in Barnett and Lewis [1994], Barnett [1976], and Beckman and Cook [1983]. We will describe a few tests here.

The *box plot rule* (Figure 9) is the simplest statistical technique that has been applied to detect univariate and multivariate anomalies in medical domain data [Laurikkala et al. 2000; Horn et al. 2001; Solberg and Lahti 2005] and turbine rotors data [Gutormsson et al. 1999]. A box plot graphically depicts the data using summary attributes such as smallest non-anomaly observation (*min*), lower quartile (Q_1), median, upper quartile (Q_3), and largest non-anomaly observation (*max*). The quantity $Q_3 - Q_1$ is called the *Inter Quartile Range* (*IQR*). The box plots also indicates the limits beyond which any observation will be treated as an anomaly. A data instance that lies more than $1.5 * IQR$ lower than Q_1 or $1.5 * IQR$ higher than Q_3 is declared as an anomaly. The region between $Q_1 - 1.5IQR$ and $Q_3 + 1.5IQR$ contains 99.3% of observations, and hence the choice of the $1.5IQR$ boundary makes the *box plot rule* equivalent to the 3σ technique for Gaussian data.

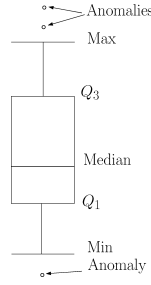


Fig. 9. A box plot for a univariate data set.

Grubb's test (also known as the *maximum normed residual test*) is used to detect anomalies in a univariate data set [Grubbs 1969; Stefansky 1972; Anscombe and Guttman 1960] under the assumption that the data is generated by a Gaussian distribution. For each test instance $\mathbf{x}$, its z score is computed as follows:

$$\mathbf{z} = \frac{|\mathbf{x} - \bar{\mathbf{x}}|}{s}, \quad (1)$$

where $\bar{\mathbf{x}}$ and s are the mean and standard deviation of the data sample, respectively. A test instance is declared to be anomalous if:

$$\mathbf{z} > \frac{N-1}{\sqrt{N}} \sqrt{\frac{t_{\alpha/(2N), N-2}^2}{N-2 + t_{\alpha/(2N), N-2}^2}}, \quad (2)$$

where N is the data size and $t_{\alpha/(2N), N-2}$ is a threshold used to declare an instance to be anomalous or normal. This threshold is the value taken by a t -distribution at a significance level of $\frac{\alpha}{2N}$. The significance level reflects the confidence associated with the threshold and indirectly controls the number of instances declared as anomalous.

A variant of the Grubb's test for multivariate data was proposed by Laurikkala et al. [2000], which uses the Mahalanobis distance of a test instance $\mathbf{x}$ to the sample mean $\bar{\mathbf{x}}$, to reduce multivariate observations to univariate scalars:

$$y^2 = (\mathbf{x} - \bar{\mathbf{x}})' S^{-1} (\mathbf{x} - \bar{\mathbf{x}}), \quad (3)$$

where S is the sample covariance matrix. The univariate *Grubb's test* is applied to y to determine if the instance $\mathbf{x}$ is anomalous or not. Several other variants of *Grubb's test* have been proposed to handle multivariate data sets [Aggarwal and Yu 2001, 2008; Laurikkala et al. 2000], graph structured data [Shekhar et al. 2001], and *Online Analytical Processing* (OLAP) data cubes [Sarawagi et al. 1998].

The student's t -test has also been applied for anomaly detection in Surace and Worden [1998] and Surace et al. [1997] to detect damages in structural beams. A normal sample, N_1 is compared with a test sample, N_2 using the t -test. If the test shows significant difference between them, it signifies the presence of an anomaly in N_2 . The multivariate version of students' t -test called the *Hotelling t^2 -test* is also used as an anomaly detection test statistic in Liu and Weng [1991] to detect anomalies in bioavailability/bioequivalence studies.

Ye and Chen [2001] use a χ^2 statistic to determine anomalies in operating system call data. The training phase assumes that the normal data has a multivariate normal

distribution. The value of the χ^2 statistic is determined as:

$$\chi^2 = \sum_{i=1}^n \frac{(X_i - E_i)^2}{E_i}, \quad (4)$$

where X_i is the observed value of the i_{th} variable, E_i is the expected value of the i_{th} variable (obtained from the training data), and n is the number of variables. A large value of χ^2 denotes that the observed sample contains anomalies.

Several other statistical anomaly detection techniques that assume that the data follows a Gaussian distribution have been proposed, but use other statistical tests, such as: *Rosner* test [Rosner 1983], *Dixon* test [Gibbons 1994], *Slippage Detection* test [Hawkins 1980], and so on.

7.1.2. Regression Model-Based. Anomaly detection using regression has been extensively investigated for time-series data [Abraham and Chuang 1989; Abraham and Box 1979; Fox 1972].

The basic regression model-based anomaly detection technique consists of two steps. In the first step, a regression model is fitted on the data. In the second step, for each test instance, the *residual* for the test instance is used to determine the anomaly score. The residual is the part of the instance which is not explained by the regression model. The magnitude of the residual can be used as the anomaly score for the test instance, though statistical tests have been proposed to determine anomalies with certain confidence [Anscombe and Guttman 1960; Beckman and Cook 1983; Hawkins 1980; Torr and Murray 1993]. Certain techniques detect the presence of anomalies in a data set by analyzing the *Akaike Information Content* (AIC) during model fitting [Kitagawa 1979; Kadota et al. 2003].

Presence of anomalies in the training data can influence the regression parameters and hence the regression model might not produce accurate results. A popular technique to handle such anomalies while fitting regression models is called *robust regression* [Rousseeuw and Leroy 1987]: estimation of regression parameters while accommodating anomalies. The authors argue that the robust regression techniques not only hide the anomalies, but can also detect the anomalies, because the anomalies tend to have larger residuals from the robust fit. A similar robust anomaly detection approach has been applied in *Autoregressive Integrated Moving Average* (ARIMA) models [Bianco et al. 2001; Chen et al. 2005].

Variants of the basic regression models-based technique have been proposed to handle multivariate time-series data. Tsay et al. [2000] discuss the additional complexity in multivariate time-series over the univariate time-series and come up with statistics that can be applied to detect anomalies in multivariate ARIMA models. This is a generalization of statistics proposed earlier by Fox [1972].

Another variant that detects anomalies in multivariate time-series data generated by an *Autoregressive Moving Average* (ARMA) model, was proposed by Galeano et al. [2004]. In this technique the authors transform the multivariate time-series to a univariate time-series by linearly combining the components of the multivariate time-series. The interesting linear combinations (projections in 1-d space) are obtained using a *projection pursuit* technique [Huber 1985] that maximizes the *Kurtosis* coefficient (a measure for the degree of peakedness/flatness in the variable distribution) of the time-series data. The anomaly detection in each projection is done by using univariate test statistics, as proposed by Fox [1972].

7.1.3. Mixture of Parametric Distributions-Based. Such techniques use a mixture of parametric statistical distributions to model the data. Techniques in this category can be grouped into two subcategories. The first subcategory of techniques model the normal instances and anomalies as separate parametric distributions, while the second subcategory of techniques model only the normal instances as a mixture of parametric distributions.

For the first subcategory of techniques, the testing phase involves determining which distribution—normal or anomalous—the test instance belongs to. Abraham and Box [1979] assume that the normal data is generated from a Gaussian distribution ($N(0, \sigma^2)$) and the anomalies are also generated from a Gaussian distribution with same mean but with larger variance, $N(0, k^2 \sigma^2)$. A test instance is tested using the Grubb's test on both distributions, and accordingly labeled as normal or anomalous. Similar techniques have been proposed in Lauer [2001], Eskin [2000], Abraham and Box [1979], Box and Tiao [1968], and Agarwal [2005]. Eskin [2000] uses the *Expectation Maximization* (EM) algorithm to develop a mixture of models for the two classes, assuming that each data point is an anomaly with apriori probability λ , and normal with apriori probability $1 - \lambda$. Thus, if $\mathbf{D}$ represents the actual probability distribution of the entire data, and $\mathbf{M}$ and $\mathbf{A}$ represent the distributions of the normal and anomalous data respectively, then $\mathbf{D} = \lambda \mathbf{A} + (1 - \lambda) \mathbf{M}$. $\mathbf{M}$ is learned using any distribution estimation technique, while $\mathbf{A}$ is assumed to be uniform. Initially all points are considered to be in $\mathbf{M}$. The anomaly score is assigned to a point based on how much the distributions change if that point is removed from $\mathbf{M}$ and added to $\mathbf{A}$.

The second subcategory of techniques model the normal instances as a mixture of parametric distributions. A test instance that does not belong to any of the learned models is declared to be an anomaly. Gaussian mixture models have been mostly used for such techniques Agarwal [2006], and have been used to detect strains in airframe data [Hickinbotham and Austin 2000a; Hollier and Austin 2002], to detect anomalies in mammographic image analysis [Spence et al. 2001; Tarassenko 1995], and for network intrusion detection [Yamanishi and ichi Takeuchi 2001; Yamanishi et al. 2004]. Similar techniques have been applied to detecting anomalies in biomedical signal data [Roberts and Tarassenko 1994; Roberts 1999, 2002], where *extreme value statistics*² are used to determine if a test point is an anomaly with respect to the learned mixture of models or not. Byers and Raftery [1998] use a mixture of Poisson distributions to model the normal data and then detect anomalies.

7.2. Nonparametric Techniques

The anomaly detection techniques in this category use nonparametric statistical models, such that the model structure is not defined a priori, but is instead determined from given data. Such techniques typically make fewer assumptions regarding the data, such as smoothness of density, when compared to parametric techniques.

7.2.1. Histogram-Based. The simplest nonparametric statistical technique is to use histograms to maintain a profile of the normal data. Such techniques are also referred to as frequency-based or counting-based. Histogram based techniques are particularly

²*Extreme Value Theory* (EVT) [Pickands 1975] is a similar concept as anomaly detection, and deals with extreme deviations of a probability distribution. EVT has been applied to risk management [McNeil 1999] as a method for modeling and measuring extreme risks. The key difference between extreme values and statistical anomalies is that extreme values are known to occur at the extremities of a probability distribution, while anomalies are more general. Anomalies can also be generated from a different distribution altogether.

popular in the intrusion detection community [Eskin 2000; Eskin et al. 2001; Denning 1987] and fraud detection [Fawcett and Provost 1999], since the behavior of the data is governed by certain profiles (user or software or system) that can be efficiently captured using the histogram model.

A basic histogram-based anomaly detection technique for univariate data consists of two steps. The first step involves building a histogram based on the different values taken by that feature in the training data. In the second step, the technique checks if a test instance falls in any one of the bins of the histogram. If it does, the test instance is normal, otherwise it is anomalous. A variant of the basic histogram-based technique is to assign an anomaly score to each test instance based on the height (frequency) of the bin in which it falls.

The size of the bin used when building the histogram is key for anomaly detection. If the bins are small, many normal test instances will fall in empty or rare bins, resulting in a high false alarm rate. If the bins are large, many anomalous test instances will fall in frequent bins, resulting in a high false negative rate. Thus a key challenge for histogram-based techniques is to determine an optimal size of the bins to construct the histogram that maintains a low false alarm rate and a low false negative rate.

Histogram-based techniques require normal data to build the histograms [Anderson et al. 1994; Javitz and Valdes 1991; Helman and Bhangoo 1997]. Some techniques even construct histograms for the anomalies [Dasgupta and Nino 2000], if labeled anomalous instances are available.

For multivariate data, a basic technique is to construct attribute-wise histograms. During testing, for each test instance, the anomaly score for each attribute value of the test instance is calculated as the height of the bin that contains the attribute value. The per-attribute anomaly scores are aggregated to obtain an overall anomaly score for the test instance.

The basic histogram-based technique for multivariate data has been applied to system call intrusion detection Endler [1998], network intrusion detection [Ho et al. 1999; Yamanishi and ichi Takeuchi 2001; Yamanishi et al. 2004], fraud detection [Fawcett and Provost 1999], damage detection in structures [Manson 2002; Manson et al. 2001, 2000], detecting Web-based attacks [Kruegel and Vigna 2003; Kruegel et al. 2002], and anomalous topic detection in text data [Allan et al. 1998]. A variant of the simple technique is used in *Packet Header Anomaly Detection* (PHAD) and *Application Layer Anomaly Detection* (ALAD) [Mahoney and Chan 2002], applied to network intrusion detection.

The SRI International's real-time *Network Intrusion Detection System* (NIDES) [Anderson et al. 1994; Anderson et al. 1995; Porras and Neumann 1997] has a subsystem that maintains long-term statistical profiles to capture the normal behavior of a computer system [Javitz and Valdes 1991]. The authors propose a *Q statistic* to compare a long-term profile with a short term profile (observation). The statistic is used to determine another measure called *S statistic*, which reflects the extent to which the behavior in a particular feature is anomalous with respect to the historical profile. The feature-wise *S statistics* are combined to get a single value called *IS statistic*, which determines if a test instance is anomalous or not. A variant has been proposed by Sargor [1998] for anomaly detection in link-state routing protocols.

7.2.2. Kernel Function-Based. A non-parametric technique for probability density estimation is *parzen windows estimation* [Parzen 1962]. This involves using kernel functions to approximate the actual density. Anomaly detection techniques based on kernel functions are similar to the parametric methods described earlier. The only difference is the density estimation technique used. Desforges et al. [1998] proposed a semi-supervised statistical technique to detect anomalies, which uses kernel functions

to estimate the probability distribution function (*pdf*) for the normal instances. A new instance, which lies in the low probability area of this *pdf* is declared to be anomalous.

Similar application of parzen windows is proposed for network intrusion detection [Chow and Yeung 2002], for novelty detection in oil flow data [Bishop 1994], and for mammographic image analysis [Tarassenko 1995].

Computational Complexity. The computational complexity of statistical anomaly detection techniques depends on the nature of the statistical model that is required to be fitted on the data. Fitting single parametric distributions from the exponential family, for example, Gaussian, Poisson, Multinomial, and so on, is typically linear in data size as well as number of attributes. Fitting complex distributions (such as mixture models, HMM, etc.) using iterative estimation techniques such as *Expectation Maximization* (EM), are also typically linear per iteration, though they might be slow in converging depending on the problem and/or convergence criterion. Kernel-based techniques can potentially have quadratic time complexity in terms of the data size.

Advantages and Disadvantages of Statistical Techniques. The advantages of statistical techniques are:

- (1) If the assumptions regarding the underlying data distribution hold true, statistical techniques provide a statistically justifiable solution for anomaly detection.
- (2) The anomaly score provided by a statistical technique is associated with a confidence interval, which can be used as additional information while making a decision regarding any test instance.
- (3) If the distribution estimation step is robust to anomalies in data, statistical techniques can operate in a unsupervised setting without any need for labeled training data.

The disadvantages of statistical techniques are:

- (1) The key disadvantage of statistical techniques is that they rely on the assumption that the data is generated from a particular distribution. This assumption often does not hold true, especially for high dimensional real data sets.
- (2) Even when the statistical assumption can be reasonably justified, there are several hypothesis test statistics that can be applied to detect anomalies; choosing the best statistic is often not a straightforward task [Motulsky 1995]. In particular, constructing hypothesis tests for complex distributions that are required to fit high dimensional data sets is nontrivial.
- (3) Histogram-based techniques are relatively simple to implement, but a key shortcoming of such techniques for multivariate data is that they are not able to capture the interactions between different attributes. An anomaly might have attribute values that are individually very frequent, but whose combination is very rare, however an attribute-wise histogram-based technique would not be able to detect such anomalies.

8. INFORMATION THEORETIC ANOMALY DETECTION TECHNIQUES

Information theoretic techniques analyze the *information content* of a data set using different information theoretic measures such as *Kolomogorov Complexity*, *entropy*, *relative entropy*, and so on. Such techniques are based on the following key assumption:

Assumption. Anomalies in data induce irregularities in the information content of the data set.

Let $\mathcal{C}(D)$ denote the complexity of a given data set, D . A basic information theoretic technique can be described as follows. Given a data set D , find the minimal subset of instances, I , such that $\mathcal{C}(D) - \mathcal{C}(D - I)$ is maximum. All instances in the subset thus obtained, are deemed as anomalous. The problem addressed by this basic technique is to find a *Pareto*-optimal solution, which does not have a single optimum, since there are two different objectives that need to be optimized.

In this technique, the complexity of a data set ($\mathcal{C}$) can be measured in different ways. *Kolomogorov complexity* [Li and Vitanyi 1993] has been used by several techniques [Arning et al. 1996; Keogh et al. 2004]. Arning et al. [1996] use the size of the regular expression to measure the *Kolomogorov Complexity* of data (represented as a string) for anomaly detection. Keogh et al. [2004] use the size of the compressed data file (using any standard compression algorithm), as a measure of the data set's *Kolomogorov Complexity*. Other information theoretic measures such as entropy, relative uncertainty, and so on, have also been used to measure the complexity of a categorical data set [Lee and Xiang 2001; He et al. 2005, 2006; Ando 2007].

This basic technique involves dual optimization to minimize the subset size while maximizing the reduction in the complexity of the data set. Thus an exhaustive approach in which every possible subset of the data set is considered would run in exponential time. Several techniques have been proposed that perform approximate search for the most anomalous subset. He et al. [2006] use an approximate algorithm called *Local Search Algorithm* (LSA) [He et al. 2005] to approximately determine such a subset in a linear fashion, using entropy as the complexity measure. A similar technique that uses an information bottleneck measure was proposed by Ando [2007].

Information theoretic techniques have also been used in data sets in which data instances are naturally ordered, for example, sequential data and spatial data. In such cases, the data is broken into substructures (segments for sequences, subgraphs for graphs, etc.), and the anomaly detection technique finds the substructure, I , such that $\mathcal{C}(D) - \mathcal{C}(D - I)$ is maximum. This technique has been applied to sequences [Lin et al. 2005; Chakrabarti et al. 1998; Arning et al. 1996], graph data [Noble and Cook 2003], and spatial data [Lin and Brown 2003]. A key challenge of such techniques is to find the optimal size of the substructure that would result in detecting anomalies.

Computational Complexity. As mentioned earlier, the basic information theoretic anomaly detection technique has exponential time complexity, though approximate techniques have been proposed that have linear time complexity.

Advantages and Disadvantages of Information Theoretic Techniques. The advantages of information theoretic techniques are as follows:

- (1) They can operate in an unsupervised setting.
- (2) They do not make any assumptions about the underlying statistical distribution for the data.

The disadvantages of information theoretic techniques are as follows:

- (1) The performance of such techniques is highly dependent on the choice of the information theoretic measure. Often, such measures can detect the presence of anomalies only when there is a significantly large number of anomalies present in the data.
- (2) Information theoretic techniques applied to sequences, and spatial data sets rely on the size of the substructure, which is often nontrivial to obtain.
- (3) It is difficult to associate an anomaly score with a test instance using an information theoretic technique.

9. SPECTRAL ANOMALY DETECTION TECHNIQUES

Spectral techniques try to find an approximation of the data using a combination of attributes that capture the bulk of the variability in the data. Such techniques are based on the following key assumption:

Assumption. Data can be embedded into a lower dimensional subspace in which normal instances and anomalies appear significantly different.

Thus the general approach adopted by spectral anomaly detection techniques is to determine such subspaces (embeddings, projections, etc.) in which the anomalous instances can be easily identified [Agovic et al. 2007]. Such techniques can work in an unsupervised as well as a semisupervised setting.

Several techniques use *Principal Component Analysis* (PCA) [Jolliffe 2002] for projecting data into a lower dimensional space. One such technique [Parra et al. 1996] analyzes the projection of each data instance along the principal components with low variance. A normal instance that satisfies the correlation structure of the data will have a low value for such projections while an anomalous instance that deviates from the correlation structure will have a large value. Dutta et al. [2007] adopt this approach to detect anomalies in astronomy catalogs.

Ide and Kashima [2004] propose a spectral technique to detect anomalies in a time series of graphs. Each graph is represented as an adjacency matrix for a given time. At every time instance, the principle component of the matrix is chosen as the *activity vector* for the given graph. The time-series of the activity vectors is considered as a matrix and the principal left singular vector is obtained to capture the normal dependencies over time in the data. For a new test graph, the angle between its activity vector and the principal left singular vector obtained from the previous graphs is computed and used to determine the anomaly score of the test graph. In a similar approach, Sun et al. [2007] propose an anomaly detection technique on a sequence of graphs by performing *Compact Matrix Decomposition* (CMD) on the adjacency matrix for each graph and thus obtaining an approximation of the original matrix. For each graph in the sequence, the authors perform CMD and compute the approximation error between the original adjacency matrix and the approximate matrix. The authors construct a time-series of the approximation errors and detect anomalies in the time-series of errors; the graph corresponding to anomalous approximation error is declared to be anomalous.

Shyu et al. [2003] present an anomaly detection technique where the authors perform robust PCA [Huber 1974] to estimate the principal components from the covariance matrix of the normal training data. The testing phase involves comparing each point with the components and assigning an anomaly score based on the point's distance from the principal components. Thus if the projection of x on the principal components is $y_1, y_2, \dots, y_p$ and the corresponding eigen-values are $\lambda_1, \lambda_2, \dots, \lambda_p$, then

$$\sum_{i=1}^q \frac{y_i^2}{\lambda_i} = \frac{y_1^2}{\lambda_1} + \frac{y_2^2}{\lambda_2} + \dots + \frac{y_q^2}{\lambda_q}, q \leq p \quad (5)$$

has a chi-square distribution [Hawkins 1974]. Using this result, the authors propose that, for a given significance level α , observation x is an anomaly if

$$\sum_{i=1}^q \frac{y_i^2}{\lambda_i} > \chi_q^2(\alpha). \quad (6)$$

It can be shown that the quantity calculated in Equation 5 is equal to the Mahalanobis distance of the instance x from the sample mean (See Equation 3) when $q = p$ [Shyu et al. 2003]. Thus the robust PCA-based technique is the same as a statistical technique discussed in Section 7.1.1 in a smaller subspace.

The robust PCA-based technique has been applied to the network intrusion detection domain [Shyu et al. 2003; Lakhina et al. 2005; Thottan and Ji 2003] and for detecting anomalies in space craft components [Fujimaki et al. 2005].

Computational Complexity. Standard PCA based techniques are typically linear in data size but often quadratic in the number of dimensions. Nonlinear techniques can improve the time complexity to be linear in the number of dimensions but polynomial in the number of principal components [Gunter et al. 2007]. Techniques that perform SVD on the data are typically quadratic in data size.

Advantages and Disadvantages of Spectral Techniques. The advantages of spectral anomaly detection techniques are as follows:

- (1) Spectral techniques automatically perform dimensionality reduction and hence are suitable for handling high dimensional data sets. Moreover, they can also be used as a preprocessing step followed by application of any existing anomaly detection technique in the transformed space.
- (2) Spectral techniques can be used in an unsupervised setting.

The disadvantages of spectral anomaly detection techniques are as follows:

- (1) Spectral techniques are useful only if the normal and anomalous instances are separable in the lower dimensional embedding of the data.
- (2) Spectral techniques typically have high computational complexity.

10. HANDLING CONTEXTUAL ANOMALIES

The anomaly detection techniques discussed in the previous sections primarily focus on detecting point anomalies. In this section, we will discuss anomaly detection techniques that handle contextual anomalies.

As discussed in Section 2.2.2, contextual anomalies require that the data has a set of *contextual attributes* (to define a context), and a set of *behavioral attributes* (to detect anomalies within a context). Song et al. [2007] use the terms *environmental* and *indicator* attributes, which are analogous to our terminology. Some of the ways in which contextual attributes can be defined are:

- (1) *Spatial.* The data has spatial attributes, which define the location of a data instance and hence a spatial neighborhood. A number of context-based anomaly detection techniques [Lu et al. 2003; Shekhar et al. 2001; Kou et al. 2006; Sun and Chawla 2004] have been proposed for data with spatial data.
- (2) *Graphs:* The edges that connect nodes (data instances) define the neighborhood for each node. Contextual anomaly detection techniques have been applied to graph-based data by Sun et al. [2005].
- (3) *Sequential:* The data is sequential: the contextual attribute of a data instance is its position in the sequence. Time-series data has been extensively explored in the contextual anomaly detection category [Abraham and Chuang 1989; Abraham and Box 1979; Rousseeuw and Leroy 1987; Bianco et al. 2001; Fox 1972; Salvador and Chan 2003; Tsay et al. 2000; Galeano et al. 2004; Zeevi et al. 1997].

Another form of sequential data for which anomaly detection techniques have been developed is event data, in which each event has a timestamp (such as operating

system call data or Web data [Ilgun et al. 1995; Vilalta and Ma 2002; Weiss and Hirsh 1998; Smyth 1994]). The difference between time series data and event sequences is that for the latter, the inter-arrival time between consecutive events is uneven.

- (4) *Profile*: Often times the data might not have an explicit spatial or sequential structure, but can still be segmented or clustered into components using a set of contextual attributes. These attributes are typically used to profile and group users in *activity monitoring* systems, such as cell-phone fraud detection [Fawcett and Provost 1999; Teng et al. 1990], CRM databases [He et al. 2004b] and credit-card fraud detection [Bolton and Hand 1999]. The users are then analyzed within their group for anomalies.

In comparison to the rich literature on point anomaly detection techniques, the research on contextual anomaly detection has been limited. Broadly, such techniques can be classified in two categories. The first category of techniques reduces a contextual anomaly detection problem to a point anomaly detection problem while the second category of techniques model the structure in the data and use the model to detect anomalies.

10.1. Reduction to Point Anomaly Detection Problem

Since contextual anomalies are individual data instances (like point anomalies), but are anomalous only with respect to a context, one approach is to apply a known point anomaly detection technique within a context.

A generic reduction-based technique consists of two steps. First, identify a context for each test instance using the contextual attributes. Second, compute an anomaly score for the test instance within the context using a known point anomaly detection technique.

An example of the generic reduction-based technique has been proposed for the scenario where identifying the context is not straightforward [Song et al. 2007]. The authors assume that the attributes are already partitioned into *contextual* and *behavioral* attributes. Thus each data instance d can be represented as $[x, y]$. The contextual data is partitioned using a mixture of Gaussian models, say U . The behavioral data is also partitioned using another mixture of Gaussian models, say V . A mapping function, $p(V_j|U_i)$ is also learned. This mapping indicates the probability of the indicator part of a data point y to be generated from a mixture component V_j , when the environmental part x is generated by U_i . Thus for a given test instance $d = [x, y]$, the anomaly score is given by:

$$Anomaly\ Score = \sum_{i=1}^{n_U} p(x \in U_i) \sum_{j=1}^{n_V} p(y \in V_j) p(V_j|U_i),$$

where n_U is the number of mixture components in U and n_V is the number of mixture components in V . $p(x \in U_i)$ indicates the probability of a sample point x to be generated from the mixture component U_i while $p(y \in U_j)$ indicates the probability of a sample point y to be generated from the mixture component V_j .

Another example of the generic technique is applied to cell phone fraud detection [Fawcett and Provost 1999]. The data in this case consists of cell phone usage records. One of the attributes in the data is the cell phone user, which is used as the contextual attribute. The activity of each user is then monitored to detect anomalies using other attributes. A similar technique is adopted for computer security [Teng et al. 1990], where the contextual attributes are: *user id*, *time of the day*. The remaining attributes are compared with existing rules representing normal behavior to detect anomalies. *Peer group analysis* [Bolton and Hand 1999] is another similar technique where users

are grouped together as *peers* and analyzed within a group for fraud. He et al. [2004b] propose the concept of *class anomaly detection*, which is essentially segmenting the data using the class labels, and then applying a known clustering-based anomaly detection technique [He et al. 2002] to detect anomalies within this subset.

For spatial data, neighborhoods are intuitive and straightforward to detect [Ng and Han 1994] by using the location coordinates. Graph-based anomaly detection [Shekhar et al. 2001; Lu et al. 2003; Kou et al. 2006] use Grubb's score [Grubbs 1969] or similar statistical point anomaly detection techniques to detect anomalies within a spatial neighborhood. Sun and Chawla [2004] use a distance-based measure called *SLOM* (Spatial Local Outlier Measure [Sun and Chawla 2006]) to detect spatial anomalies within a neighborhood.

Another example of the generic technique applied to time-series data is proposed by Basu and Meckesheimer [2007]. For a given instance in a time-series the authors compare the observed value to the median of the neighborhood values. A transformation technique for time-series data has been proposed by using phase spaces [Ma and Perkins 2003b]. This technique converts a time-series into a set of vectors by unfolding the time-series into a phase space using a time delay embedding process. The temporal relations at any time instance are embedded in the phase vector for that instance. The authors use this technique to transform a time-series into feature space and then use one-class SVMs to detect anomalies. Each anomaly can be translated to a value at certain time instance in the original time-series.

10.2. Utilizing the Structure in Data

In several scenarios, breaking up data into contexts is not straightforward. This is typically true for time-series data and event sequence data. In such cases, time-series modeling and sequence modeling techniques are extended to detect contextual anomalies in the data.

A generic technique in this category can be described as follows. A model is learned from the training data, which can predict the expected behavior with respect to a given context. If the expected behavior is significantly different from the observed behavior, an anomaly is declared. A simple example of this generic technique is regression in which the contextual attributes can be used to predict the behavioral attribute by fitting a regression line on the data.

For time series data, several regression based techniques for time-series modeling such as robust regression [Rousseeuw and Leroy 1987], auto-regressive models [Fox 1972], ARMA models [Abraham and Chuang 1989; Abraham and Box 1979; Galeano et al. 2004; Zeevi et al. 1997], and ARIMA models [Bianco et al. 2001; Tsay et al. 2000], have been developed for contextual anomaly detection. Regression-based techniques have been extended to detect contextual anomalies in a set of coevolving sequences by modeling the regression as well as correlation between the sequences [Yi et al. 2000].

One of the earliest works in time-series anomaly detection was proposed by Fox [1972], where a time-series was modeled as a stationary autoregressive process. Any observation is tested to be anomalous by comparing it with the covariance matrix of the autoregressive process. If the observation falls outside the modeled error for the process, it is declared to be an anomaly. An extension to this technique is made by using *Support Vector Regression* to estimate the regression parameters and then using the learned model to detect novelties in the data [Ma and Perkins 2003a].

A technique to detect a single anomaly (discord) in a sequence of alphabets was proposed by Keogh et al. [2004]. The technique adopts a divide and conquer approach. The sequence is divided into two parts and the *Kolmogorov Complexity* is calculated for

each. The one with higher complexity contains the anomaly. The sequence is recursively divided until they are left with a single event which is declared to be the anomaly in the sequence.

Weiss and Hirsh [1998] propose a technique to detect rare events in sequential data, where they use events occurring before a particular time to predict the event occurring at that time instance. If the prediction does not match the actual event, it is declared to be rare. This idea is extended in other areas, where the authors have used Frequent Itemset Mining [Vilalta and Ma 2002], *Finite State Automaton* (FSA) [Ilgun et al. 1995; Salvador and Chan 2003] and Markov Models [Smyth 1994] to determine conditional probabilities for events based on the history of events. Marceau [2000] uses FSA to predict the next event of a sequence based on the previous n events. They apply this technique to the domain of system call intrusion detection. Hollmen and Tresp [1999] employ HMM for cell phone fraud detection. The authors use a *hierarchical regime switching call model* to model the cell phone activity of a user. The model predicts the probability of a fraud taking place for a call using the learned model. The parameter estimation is done using the EM algorithm.

A model to detect intrusions in telephone networks was proposed by Scott [2001] and for modeling Web click data by Ihler et al. [2006]. Both articles follow a technique in which they assume that the normal behavior in a time-series is generated by a nonstationary Poisson process while the anomalies are generated by a homogenous Poisson process. The transition between normal and anomalous behavior is modeled using a Markov process. The proposed techniques in all of these articles use the *Markov Chain Monte Carlo* (MCMC) estimation technique to estimate the parameters for these processes. For testing, a time series is modeled using this process and the time instances for which the anomalous behavior was active are considered as anomalies.

Bipartite graph structure in P2P networks has been used to first identify a neighborhood for any node in the graph [Sun et al. 2005], and then detecting the relevance of that node within the neighborhood. A node with a low relevance score is treated as an anomaly. The authors also propose an approximate technique where the graph is first partitioned into nonoverlapping subgraphs using graph partitioning algorithms such as METIS [Karypis and Kumar 1998]. The neighborhood of a node is then computed within its partition.

Computational Complexity. The computational complexity of the training phase in reduction-based contextual anomaly detection techniques depends on the reduction technique as well as the point anomaly detection technique used within each context. While segmenting/partitioning techniques have a fast reduction step, techniques that use clustering, or a mixture of models estimation, are relatively slower. Since the reduction simplifies the anomaly detection problem, fast point anomaly detection techniques can be used to speed up the second step. The testing phase is relatively expensive since for each test instance, its context is determined, and then an anomaly label or score is assigned using a point anomaly detection technique.

The computational complexity of the training phase in contextual anomaly detection techniques that utilize the structure in the data to build models, is typically higher than that of techniques that reduce the problem to point anomaly detection. An advantage for such techniques is the testing phase is relatively fast, since each instance is just compared to the single model and assigned an anomaly score or an anomaly label.

Advantages and Disadvantages of Contextual Anomaly Detection Techniques. The key advantage of contextual anomaly detection techniques is that they allow a natural definition of an anomaly in many real-life applications where data instances tend to be

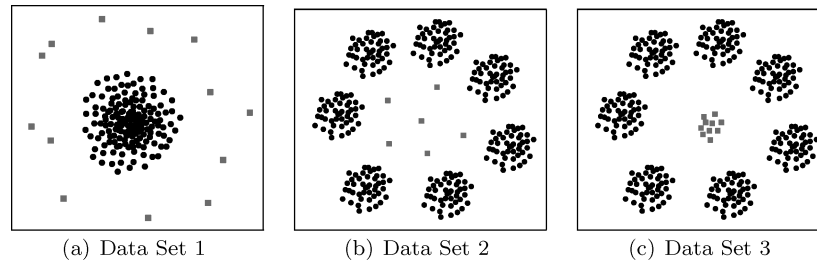


Fig. 10. 2-D data sets. Normal instances are shown as circles and anomalies are shown as squares.

similar within a context. Such techniques are able to detect anomalies that might not be detected by point anomaly detection techniques that take a global view of the data.

The disadvantage of contextual anomaly detection techniques is that they are applicable only when a context can be defined.

11. RELATIVE STRENGTHS AND WEAKNESSES OF ANOMALY DETECTION TECHNIQUES

Each of the large number of anomaly detection techniques discussed in the previous sections have their unique strengths and weaknesses. It is important to know which anomaly detection technique is best suited for a given problem. Given the complexity of the problem space, it is not feasible to provide such an understanding for every anomaly detection problem. In this section we analyze the relative strengths and weaknesses of different categories of techniques for a few simple problem settings.

For example, let us consider the following anomaly detection problem. The input is 2-D continuous data (Figure 10(a)). The normal data instances are generated from a Gaussian distribution and are located in a tight cluster in the 2-D space. The anomalies are a very few instances generated from another Gaussian distribution whose mean is very far from the first distribution. A representative training data set that contains instances from the normal data set is also available. Thus the assumptions made by techniques in Sections 4–9 hold for this data set, and hence any anomaly detection techniques belonging to these categories will detect the anomalies in such a scenario.

Now let us consider another 2-D data set (Figure 10(b)). Let the normal instances be such that they are generated by a large number of different Gaussian distributions with means arranged on a circle and very low variance. Thus the normal data will be a set of tight clusters arranged on a circle. A one-class classification-based technique might learn a circular boundary around the entire data set and hence will not be able to detect the anomalies that lie within the circle of clusters. On the other hand if each cluster was labeled as a different class, a multi-class classification-based technique might be able to learn boundaries around each cluster, and hence be able to detect the anomalies in the center. A statistical technique that uses a mixture model approach to model the data, may be able to detect the anomalies. Similarly, clustering-based and nearest neighbor based techniques will be able to detect the anomalies since they are far from all other instances. In a similar example (Figure 10(c)), if the anomalous instances form a tight cluster of significant size at the center of the circle, both clustering-based and nearest neighbor-based techniques will treat these instances as normal, thus exhibiting poor performance.

For more complex data sets, different types of techniques face different challenges. Nearest neighbor and clustering-based techniques suffer when the number of dimensions is high, because the distance measures in a high number of dimensions are not able to differentiate between normal and anomalous instances. Spectral techniques

explicitly address the high dimensionality problem by mapping data to a lower dimensional projection. But their performance is highly dependent on the assumption that the normal instances and anomalies are distinguishable in the projected space. Classification-based techniques can be a better choice in such scenario. But to be most effective, classification-based techniques require labels for both normal and anomalous instances, which are not often available. Even if the labels for both normal and anomalous instances are available, the imbalance in the distribution of the two labels often makes learning a classifier quite challenging. Semisupervised nearest neighbor and clustering techniques, that only use the normal labels, can often be more effective than the classification-based techniques. Statistical techniques, though unsupervised, are effective only when the dimensionality of data is low and statistical assumptions hold. Information theoretic techniques require a measure that is sensitive enough to detect the effects of even a single anomaly. Otherwise, such techniques can detect anomalies only when there are a significantly adequate number of anomalies.

Nearest neighbor and clustering-based techniques require distance computation between a pair of data instances. Thus, such techniques assume that the distance measure can discriminate between the anomalies and normal instances well enough. In situations where identifying a good distance measure is difficult, classification-based or statistical techniques might be a better choice.

The computational complexity of an anomaly detection technique is a key aspect, especially when the technique is applied to a real domain. While classification-based, clustering-based, and statistical techniques have expensive training times, testing is usually fast. Often this is acceptable, since models can be trained in an offline fashion while testing is required to be in real time. In contrast, techniques such as nearest neighbor-based, information theoretic, and spectral techniques, which do not have a training phase, have an expensive testing phase, which can be a limitation in a real setting.

Anomaly detection techniques typically assume that anomalies in data are rare when compared to normal instances. Though this assumption is generally true, anomalies are not always rare. For example, when dealing with worm detection in computer networks, the anomalous (worm) traffic is actually more frequent than the normal traffic. Unsupervised techniques are not suited for such bulk anomaly detection. Techniques operating in supervised or semisupervised modes can be applied to detect bulk anomalies [Sun et al. 2007; Soule et al. 2005].

12. CONCLUDING REMARKS AND FUTURE WORK

In this survey we have discussed different ways in which the problem of anomaly detection has been formulated in the literature, and have attempted to provide an overview of the huge literature on various techniques. For each category of anomaly detection techniques, we have identified a unique assumption regarding the notion of normal and anomalous data. When applying a given technique to a particular domain, these assumptions can be used as guidelines to assess the effectiveness of the technique in that domain. Ideally, a comprehensive survey on anomaly detection should allow a reader to not only understand the motivation behind using a particular anomaly detection technique, but also provide a comparative analysis of various techniques. But the current research has been done in an unstructured fashion, without relying on a unified notion of anomalies, which makes the job of providing a theoretical understanding of the anomaly detection problem very difficult. A possible future work would be to unify the assumptions made by different techniques regarding the normal and anomalous behavior into a statistical or machine learning framework. A limited attempt in this

direction is provided by Knorr and Ng [1997], where the authors show the relation between distance-based and statistical anomalies for two-dimensional data sets.

There are several promising directions for further research in anomaly detection. Contextual and collective anomaly detection techniques are beginning to find increasing applicability in several domains and there is much scope for development of new techniques in this area. The presence of data across different distributed locations has motivated the need for distributed anomaly detection techniques [Zimmermann and Mohay 2006]. While such techniques process information available at multiple sites, they often have to simultaneously protect the information present in each site, thereby requiring privacy-preserving anomaly detection techniques [Vaidya and Clifton 2004]. With the emergence of sensor networks, processing data as it arrives has become a necessity. Many techniques discussed in this survey require the entire test data before detecting anomalies. Recently, techniques have been proposed that can operate in an online fashion [Pokrajac et al. 2007]; such techniques not only assign an anomaly score to a test instance as it arrives, but also incrementally update the model. Another upcoming area where anomaly detection is finding more and more applicability is in complex systems. An example of such a system would be an aircraft system with multiple components. Anomaly detection in such systems involves modeling the interaction among various components [Bronstein et al. 2001].

ACKNOWLEDGMENTS

The authors thank Shyam Boriah and Gang Fang for extensive comments on the final draft of the article.

REFERENCES

- ABE, N., ZADROZNY, B., AND LANGFORD, J. 2006. Outlier detection by active learning. In *Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, New York, 504–509.
- ABRAHAM, B. AND BOX, G. E. P. 1979. Bayesian analysis of some outlier problems in time series. *Biometrika* 66, 2, 229–236.
- ABRAHAM, B. AND CHUANG, A. 1989. Outlier detection and time series modeling. *Technometrics* 31, 2, 241–248.
- ADDISON, J., WERMTER, S., AND MACINTYRE, J. 1999. Effectiveness of feature extraction in neural network architectures for novelty detection. In *Proceedings of the 9th International Conference on Artificial Neural Networks*. vol. 2. 976–981.
- AEYELS, D. 1991. On the dynamic behaviour of the novelty detector and the novelty filter. In *Analysis of Controlled Dynamical Systems: Progress in Systems and Control Theory*, B. Bonnard, B. Bride, J. Gauthier, and I. Kupka, Eds. vol. 8. Springer, Berlin, 1–10.
- AGARWAL, D. 2005. An empirical Bayes approach to detect anomalies in dynamic multidimensional arrays. In *Proceedings of the 5th IEEE International Conference on Data Mining*. IEEE Computer Society, 26–33.
- AGARWAL, D. 2006. Detecting anomalies in cross-classified streams: A Bayesian approach. *Knowl. Inform. Syst.* 11, 1, 29–44.
- AGGARWAL, C. 2005. On abnormality detection in spuriously populated data streams. In *Proceedings of the 5th SIAM Data Min. Conference*. 80–91.
- AGGARWAL, C. AND YU, P. 2001. Outlier detection for high dimensional data. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*. ACM Press, 37–46.
- AGGARWAL, C. C. AND YU, P. S. 2008. Outlier detection with uncertain data. In *Proceedings of the International Conference on Data Mining (SDM)*. 483–493.
- AGOVIC, A., BANERJEE, A., GANGULY, A. R., AND PROTOPOPESCU, V. 2007. Anomaly detection in transportation corridors using manifold embedding. In *Proceedings of the 1st International Workshop on Knowledge Discovery from Sensor Data*. ACM Press.
- AGRAWAL, R. AND SRIKANT, R. 1995. Mining sequential patterns. In *Proceedings of the 11th International Conference on Data Engineering*. IEEE Computer Society, 3–14.

- AGYEMANG, M., BARKER, K., AND ALHAJJ, R. 2006. A comprehensive survey of numeric and symbolic outlier mining techniques. *Intel. Data Anal.* 10, 6, 521–538.
- ALBRECHT, S., BUSCH, J., KLOPPENBURG, M., METZE, F., AND TAVAN, P. 2000. Generalized radial basis function networks for classification and novelty detection: Self-organization of optional Bayesian decision. *Neural Netw.* 13, 10, 1075–1093.
- ALESKEROV, E., FREISLEBEN, B., AND RAO, B. 1997. Cardwatch: A neural network based database mining system for credit card fraud detection. In *Proceedings of the IEEE Conference on Computational Intelligence for Financial Engineering*. 220–226.
- ALLAN, J., CARBONELL, J., DODDINGTON, G., YAMRON, J., AND YANG, Y. 1998. Topic detection and tracking pilot study. In *Proceedings of the DARPA Broadcast News Transcription and Understanding Workshop*. 194–218.
- ANDERSON, D. LUNT, T. F., JAVITZ, H., TAMARU, A., AND VALDES, A. 1995. Detecting unusual program behavior using the statistical components of NIDES. Tech. rep. SRI-CSL-95-06, Computer Science Laboratory, SRI International.
- ANDERSON, D., FRIVOLD, T., TAMARU, A., AND VALDES, A. 1994. Next-generation intrusion detection expert system (NIDES), software users manual, beta-update release. Tech. rep. SRI-CSL-95-07, Computer Science Laboratory, SRI International.
- ANDO, S. 2007. Clustering needles in a haystack: An information theoretic analysis of minority and outlier detection. In *Proceedings of the 7th International Conference on Data Mining*. 13–22.
- ANGIULLI, F. AND PIZZUTI, C. 2002. Fast outlier detection in high dimensional spaces. In *Proceedings of the 6th European Conference on Principles of Data Mining and Knowledge Discovery*. Springer-Verlag, 15–26.
- ANScombe, F. J. AND GUTTMAN, I. 1960. Rejection of outliers. *Technometrics* 2, 2, 123–147.
- ARNING, A., AGRAWAL, R., AND RAGHAVAN, P. 1996. A linear method for deviation detection in large databases. In *Proceedings of the 2nd International Conference of Knowledge Discovery and Data Mining*. 164–169.
- AUGUSTEIJN, M. AND FOLKERT, B. 2002. Neural network classification and novelty detection. *Int. J. Rem. Sens.* 23, 14, 2891–2902.
- BAKAR, Z., MOHEMAD, R., AHMAD, A., AND DERIS, M. 2006. A comparative study for outlier detection techniques in data mining. *Proceedings of the IEEE Conference on Cybernetics and Intelligent Systems*. 1–6.
- BAKER, D., HOFMANN, T., MCCALLUM, A., AND YANG, Y. 1999. A hierarchical probabilistic model for novelty detection in text. In *Proceedings of the International Conference on Machine Learning*.
- BARBARA, D., COUTO, J., JAJODIA, S., AND WU, N. 2001a. Adam: A testbed for exploring the use of data mining in intrusion detection. *SIGMOD Rec.* 30, 4, 15–24.
- BARBARA, D., COUTO, J., JAJODIA, S., AND WU, N. 2001b. Detecting novel network intrusions using Bayes estimators. In *Proceedings of the 1st SIAM International Conference on Data Mining*.
- BARBARA, D., LI, Y., COUTO, J., LIN, J.-L., AND JAJODIA, S. 2003. Bootstrapping a data mining intrusion detection system. In *Proceedings of the ACM Symposium on Applied Computing*. ACM Press, 421–425.
- BARNETT, V. 1976. The ordering of multivariate data (with discussion). *J. Royal Statist. Soc. Series A* 139, 318–354.
- BARNETT, V. AND LEWIS, T. 1994. *Outliers in Statistical Data*. John Wiley.
- BARSON, P., DAVEY, N., FIELD, S. D. H., FRANK, R. J., AND McASKIE, G. 1996. The detection of fraud in mobile phone networks. *Neural Netw. World* 6, 4.
- BASU, S., BLENKO, M., AND MOONEY, R. J. 2004. A probabilistic framework for semi-supervised clustering. In *Proceedings of the 10th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 59–68.
- BASU, S. AND MECKESHEIMER, M. 2007. Automatic outlier detection for time series: an application to sensor data. *Know. Inform. Syst.* 11, 2, 137–154.
- BAY, S. D. AND SCHWABACHER, M. 2003. Mining distance-based outliers in near linear time with randomization and a simple pruning rule. In *Proceedings of the 9th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 29–38.
- BECKMAN, R. J. AND COOK, R. D. 1983. Outlier...s. *Technometrics* 25, 2, 119–149.
- BEJERANO, G. AND YONA, G. 2001. Variations on probabilistic suffix trees: statistical modeling and prediction of protein families. *Bioinformatics* 17, 1, 23–43.
- BENTLEY, J. L. 1975. Multi-dimensional binary search trees used for associative searching. *Comm. ACM* 18, 9, 509–517.
- BIANCO, A. M., BEN, M. G., MARTINEZ, E. J., AND YOHAI, V. J. 2001. Outlier detection in regression models with arima errors using robust estimates. *J. Forecast.* 20, 8, 565–579.

- BISHOP, C. 1994. Novelty detection and neural network validation. In *Proceedings of the IEEE Conference on Vision, Image and Signal Processing*. vol. 141. 217–222.
- BLENDER, R., FRAEDRICH, K., AND LUNKEIT, F. 1997. Identification of cyclone-track regimes in the north atlantic. *Quart. J. Royal Meteor. Soc.* 123, 539, 727–741.
- BOLTON, R. AND HAND, D. 1999. Unsupervised profiling methods for fraud detection. In *Proceedings of the Conference on Credit Scoring and Credit Control VII*.
- BORIAH, S., CHANDOLA, V., AND KUMAR, V. 2008. Similarity measures for categorical data: A comparative evaluation. In *Proceedings of the 8th SIAM International Conference on Data Mining*. 243–254.
- BORISYUK, R., DENHAM, M., HOPPENSTEADT, F., KAZANOVICH, Y., AND VINOGRADOVA, O. 2000. An oscillatory neural network model of sparse distributed memory and novelty detection. *Biosystems* 58, 265–272.
- BOX, G. E. P. AND TIAO, G. C. 1968. Bayesian analysis of some outlier problems. *Biometrika* 55, 1, 119–129.
- BRANCH, J., SZYMANSKI, B., GIANNELLA, C., WOLFF, R., AND KARGUPTA, H. 2006. In-network outlier detection in wireless sensor networks. In *Proceedings of the 26th IEEE International Conference on Distributed Computing Systems*.
- BRAUSE, R., LANGSDORF, T., AND HEPP, M. 1999. Neural data mining for credit card fraud detection. In *Proceedings of the IEEE International Conference on Tools with Artificial Intelligence*. 103–106.
- BREUNIG, M. M., KRIEGEL, H.-P., NG, R. T., AND SANDER, J. 1999. Optics-of: Identifying local outliers. In *Proceedings of the 3rd European Conference on Principles of Data Mining and Knowledge Discovery*. Springer-Verlag, 262–270.
- BREUNIG, M. M., KRIEGEL, H.-P., NG, R. T., AND SANDER, J. 2000. LOF: Identifying density-based local outliers. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*. ACM Press, 93–104.
- BRITO, M. R., CHAVEZ, E. L., QUIROZ, A. J., AND YUKICH, J. E. 1997. Connectivity of the mutual k-nearest-neighbor graph in clustering and outlier detection. *Statist. Prob. Lett.* 35, 1, 33–42.
- BROCKETT, P. L., XIA, X., AND DERRIG, R. A. 1998. Using Kohonen's self-organizing feature map to uncover automobile bodily injury claims fraud. *J. Risk Insur.* 65, 2, 245–274.
- BRONSTEIN, A., DAS, J., DURO, M., FRIEDRICH, R., KLEYNER, G., MUELLER, M., SINGHAL, S., AND COHEN, I. 2001. Bayesian networks for detecting anomalies in Internet-based services. In *Proceedings of the International Symposium on Integrated Network Management*.
- BROTHERTON, T. AND JOHNSON, T. 2001. Anomaly detection for advanced military aircraft using neural networks. In *Proceedings of the IEEE Aerospace Conference*.
- BROTHERTON, T., JOHNSON, T., AND CHADDERDON, G. 1998. Classification and novelty detection using linear models and a class dependent-elliptical basis function neural network. In *Proceedings of the IJCNN Conference*.
- BUDALAKOTI, S., SRIVASTAVA, A., AKELLA, R., AND TURKOV, E. 2006. Anomaly detection in large sets of high-dimensional symbol sequences. Tech. rep. NASA TM-2006-214553, NASA Ames Research Center.
- BYERS, S. D. AND RAFTERY, A. E. 1998. Nearest neighbor clutter removal for estimating features in spatial point processes. *J. Amer. Statist. Assoc.* 93, 577–584.
- BYUNGHO, H. AND SUNGZON, C. 1999. Characteristics of autoassociative MLP as a novelty detector. In *Proceedings of the IEEE International Joint Conference on Neural Networks*. Vol. 5. 3086–3091.
- CABRERA, J. B. D., LEWIS, L., AND MEHRA, R. K. 2001. Detection and classification of intrusions and faults using sequences of system calls. *SIGMOD Rec.* 30, 4, 25–34.
- CAMPBELL, C. AND BENNETT, K. 2001. A linear programming approach to novelty detection. In *Proceedings of the Conference on Advances in Neural Information Processing*. vol. 14. Cambridge Press.
- CAUDELL, T. AND NEWMAN, D. 1993. An adaptive resonance architecture to define normality and detect novelties in time series and databases. In *Proceedings of the IEEE World Congress on Neural Networks*. IEEE, 166–176.
- CHAKRABARTI, S., SARAWAGI, S., AND DOM, B. 1998. Mining surprising patterns using temporal description length. In *Proceedings of the 24th International Conference on Very Large Data Bases*. Morgan Kaufmann Publishers Inc., 606–617.
- CHANDOLA, V., BANERJEE, A., AND KUMAR, V. 2007. Anomaly detection: A survey. Tech. rep. 07-017, Computer Science Department, University of Minnesota.
- CHANDOLA, V., BORIAH, S., AND KUMAR, V. 2008. Understanding categorical similarity measures for outlier detection. Tech. rep. 08-008, University of Minnesota.
- CHANDOLA, V., EILERTSON, E., ERTÖZ, L., SIMON, G., AND KUMAR, V. 2006. Data mining for cyber security. In *Data Warehousing and Data Mining Techniques for Computer Security*, A. Singhal, Ed. Springer.

- CHATZIGIANNAKIS, V., PAPAVALASSIOU, S., GRAMMATIKOU, M., AND MAGLARIS, B. 2006. Hierarchical anomaly detection in distributed large-scale sensor networks. In *Proceedings of the 11th IEEE Symposium on Computers and Communications (ISCC)*. IEEE Computer Society, 761–767.
- CHAUDHARY, A., SZALAY, A. S., AND MOORE, A. W. 2002. Very fast outlier detection in large multidimensional data sets. In *Proceedings of the ACM SIGMOD Workshop in Research Issues in Data Mining and Knowledge Discovery (DMKD)*. ACM Press.
- CHAWLA, N. V., JAPKOWICZ, N., AND KOTCZ, A. 2004. Editorial: special issue on learning from imbalanced data sets. *SIGKDD Explor.* 6, 1, 1–6.
- CHEN, D., SHAO, X., HU, B., AND SU, Q. 2005. Simultaneous wavelength selection and outlier detection in multivariate regression of near-infrared spectra. *Anal. Sci.* 21, 2, 161–167.
- CHIU, A. AND CHEE FU, A. W. 2003. Enhancements on local outlier detection. In *Proceedings of the 7th International Database Engineering and Applications Symposium*. 298–307.
- CHOW, C. AND YEUNG, D.-Y. 2002. Parzen-window network intrusion detectors. In *Proceedings of the 16th International Conference on Pattern Recognition*. vol. 4. IEEE Computer Society, 40385.
- COX, K. C., EICK, S. G., WILLS, G. J., AND BRACHMAN, R. J. 1997. Visual data mining: Recognizing telephone calling fraud. *J. Data Min. Knowl. Disc.* 1, 2, 225–231.
- CROOK, P. AND HAYES, G. 2001. A robot implementation of a biologically inspired method for novelty detection. In *Proceedings of the Towards Intelligent Mobile Robots Conference*.
- CROOK, P. A., MARSLAND, S., HAYES, G., AND NEHMZOW, U. 2002. A tale of two filters: Online novelty detection. In *Proceedings of the International Conference on Robotics and Automation*. 3894–3899.
- CUN, Y. L., BOSER, B., DENKER, J. S., HOWARD, R. E., HUBBARD, W., JACKEL, L. D., AND HENDERSON, D. 1990. Handwritten digit recognition with a back-propagation network. In *Advances in Neural Information Processing Systems*. 396–404. Morgan Kaufmann.
- DAS, K. AND SCHNEIDER, J. 2007. Detecting anomalous records in categorical datasets. In *Proceedings of the 13th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press.
- DASGUPTA, D. AND MAJUMDAR, N. 2002. Anomaly detection in multidimensional data using negative selection algorithm. In *Proceedings of the IEEE Conference on Evolutionary Computation*. 1039–1044.
- DASGUPTA, D. AND NINO, F. 2000. A comparison of negative and positive selection algorithms in novel pattern detection. In *Proceedings of the IEEE International Conference on Systems, Man, and Cybernetics*. vol. 1. 125–130.
- DAVY, M. AND GODSILL, S. 2002. Detection of abrupt spectral changes using support vector machines, an application to audio signal segmentation. In *Proceedings of the IEEE International Conference on Acoustics, Speech, and Signal Processing*.
- DEBAR, H., DACIER, M., NASSEHI, M., AND WESPI, A. 1998. Fixed vs. variable-length patterns for detecting suspicious process behavior. In *Proceedings of the 5th European Symposium on Research in Computer Security*. Springer-Verlag, 1–15.
- DENNING, D. E. 1987. An intrusion detection model. *IEEE Trans. Softw. Eng.* 13, 2, 222–232.
- DESFORGES, M., JACOB, P., AND COOPER, J. 1998. Applications of probability density estimation to the detection of abnormal conditions in engineering. In *Proceedings of the Institute of the Mechanical Engineers*. vol. 212. 687–703.
- DIAZ, I. AND HOLLMER, J. 2002. Residual generation and visualization for understanding novel process conditions. In *Proceedings of the IEEE International Joint Conference on Neural Networks*. IEEE, 2070–2075.
- DIEHL, C. AND HAMPSHIRE, J. 2002. Real-time object classification and novelty detection for collaborative video surveillance. In *Proceedings of the IEEE International Joint Conference on Neural Networks*. IEEE.
- DONOHU, S. 2004. Early detection of insider trading in option markets. In *Proceedings of the 10th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 420–429.
- DORRONSORO, J. R., GINEL, F., SANCHEZ, C., AND CRUZ, C. S. 1997. Neural fraud detection in credit card operations. *IEEE Trans. Neural Netw.* 8, 4, 827–834.
- DU, W., FANG, L., AND PENG, N. 2006. Lad: Localization anomaly detection for wireless sensor networks. *J. Paral. Distrib. Comput.* 66, 7, 874–886.
- DUDA, R. O., HART, P. E., AND STORK, D. G. 2000. *Pattern Classification 2nd Ed.* Wiley-Interscience.
- DUTTA, H., GIANNELLA, C., BORNE, K., AND KARGUPTA, H. 2007. Distributed top-k outlier detection in astronomy catalogs using the DEMAC system. In *Proceedings of the 7th SIAM International Conference on Data Mining*.
- EDGEWORTH, F. Y. 1887. On discordant observations. *Philosoph. Mag.* 23, 5, 364–375.

- EMAMIAN, V., KAVEH, M., AND TEWFIK, A. 2000. Robust clustering of acoustic emission signals using the Kohonen network. In *Proceedings of the IEEE International Conference of Acoustics, Speech and Signal Processing*. IEEE Computer Society.
- ENDLER, D. 1998. Intrusion detection: Applying machine learning to solaris audit data. In *Proceedings of the 14th Annual Computer Security Applications Conference*. IEEE Computer Society, 268.
- ERTOZ, L., EILERTSON, E., LAZAREVIC, A., TAN, P.-N., KUMAR, V., SRIVASTAVA, J., AND DOKAS, P. 2004. MINDS—Minnesota Intrusion Detection System. In *Data Mining—Next Generation Challenges and Future Directions*. MIT Press.
- ERTÖZ, L., STEINBACH, M., AND KUMAR, V. 2003. Finding topics in collections of documents: A shared nearest neighbor approach. In *Clustering and Information Retrieval*. 83–104.
- ESCALANTE, H. J. 2005. A comparison of outlier detection algorithms for machine learning. In *Proceedings of the International Conference on Communications in Computing*.
- ESKIN, E. 2000. Anomaly detection over noisy data using learned probability distributions. In *Proceedings of the 17th International Conference on Machine Learning*. Morgan Kaufmann Publishers Inc., 255–262.
- ESKIN, E., ARNOLD, A., PRERAU, M., PORTNOY, L., AND STOLFO, S. 2002. A geometric framework for unsupervised anomaly detection. In *Proceedings of the Conference on Applications of Data Mining in Computer Security*. Kluwer Academics, 78–100.
- ESKIN, E., LEE, W., AND STOLFO, S. 2001. Modeling system call for intrusion detection using dynamic window sizes. In *Proceedings of DARPA Information Survivability Conference and Exposition (DISCEX)*.
- ESTER, M., KRIEGEL, H.-P., SANDER, J., AND XU, X. 1996. A density-based algorithm for discovering clusters in large spatial databases with noise. In *Proceedings of the 2nd International Conference on Knowledge Discovery and Data Mining*, E. Simoudis, J. Han, and U. Fayyad, Eds. AAAI Press, 226–231.
- FAN, W., MILLER, M., STOLFO, S. J., LEE, W., AND CHAN, P. K. 2001. Using artificial anomalies to detect unknown and known network intrusions. In *Proceedings of the IEEE International Conference on Data Mining*. IEEE Computer Society, 123–130.
- FAWCETT, T. AND PROVOST, F. 1999. Activity monitoring: noticing interesting changes in behavior. In *Proceedings of the 5th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 53–62.
- FORREST, S., D'HAESELEER, P., AND HELMAN, P. 1996a. An immunological approach to change detection: Algorithms, analysis and implications. In *Proceedings of the IEEE Symposium on Security and Privacy*. IEEE Computer Society, 110.
- FORREST, S., ESPONDA, F., AND HELMAN, P. 2004. A formal framework for positive and negative detection schemes. In *IEEE Trans. Syst. Man Cybernetics, Part B*. IEEE, 357–373.
- FORREST, S., HOFMEYR, S. A., SOMAYAJI, A., AND LONGSTAFF, T. A. 1996b. A sense of self for unix processes. In *Proceedings of the IEEE ISRSP*. 120–128.
- FORREST, S., PERELSON, A. S., ALLEN, L., AND CHERUKURI, R. 1994. Self-nonsel self discrimination in a computer. In *Proceedings of the IEEE Symposium on Security and Privacy*. IEEE Computer Society, 202.
- FORREST, S., WARRENDER, C., AND PEARLMUTTER, B. 1999. Detecting intrusions using system calls: Alternate data models. In *Proceedings of the IEEE ISRSP*. IEEE Computer Society, 133–145.
- FOX, A. J. 1972. Outliers in time series. *J. Royal Statist. Soc. Series B* 34, 3, 350–363.
- FUJIMAKI, R., YAIRI, T., AND MACHIDA, K. 2005. An approach to spacecraft anomaly detection problem using kernel feature space. In *Proceedings of the 11th ACM SIGKDD International Conference on Knowledge Discovery in Data Mining*. ACM Press, 401–410.
- GALEANO, P., PEA, D., AND TSAY, R. S. 2004. Outlier detection in multivariate time series via projection pursuit. Statistics and econometrics working articles ws044211, Departamento de Estadística y Econometría, Universidad Carlos III.
- GHOSH, A. K., SCHWARTZBARD, A., AND SCHATZ, M. 1999a. Learning program behavior profiles for intrusion detection. In *Proceedings of the 1st USENIX Workshop on Intrusion Detection and Network Monitoring*. 51–62.
- GHOSH, A. K., SCHWARTZBARD, A., AND SCHATZ, M. 1999b. Using program behavior profiles for intrusion detection. In *Proceedings of the SANS 3rd Conference and Workshop on Intrusion Detection and Response*.
- GHOSH, A. K., WANKEN, J., AND CHARRON, F. 1998. Detecting anomalous and unknown intrusions against programs. In *Proceedings of the 14th Annual Computer Security Applications Conference*. IEEE Computer Society, 259.
- GHOSH, S. AND REILLY, D. L. 1994. Credit card fraud detection with a neural-network. In *Proceedings of the 27th Annual Hawaii International Conference on System Science*. vol. 3.

- GHOTING, A., PARTHASARATHY, S., AND OTEY, M. 2006. Fast mining of distance-based outliers in high dimensional datasets. In *Proceedings of the SIAM International Conference on Data Mining*.
- GIBBONS, R. D. 1994. *Statistical Methods for Groundwater Monitoring*. John Wiley & Sons, Inc.
- GOLDBERGER, A. L., AMARAL, L. A. N., GLASS, L., HAUSDORFF, J. M., IVANOV, P. C., MARK, R. G., MIETUS, J. E., MOODY, G. B., PENG, C.-K., AND STANLEY, H. E. 2000. PhysioBank, PhysioToolkit, and PhysioNet: Components of a new research resource for complex physiologic signals. *Circulation* 101, 23, e215–e220. Circulation Electronic Pages: <http://circ.ahajournals.org/cgi/content/full/101/23/e215>.
- GONZALEZ, F. A. AND DASGUPTA, D. 2003. Anomaly detection using real-valued negative selection. *Genetic Program. Evol. Mach.* 4, 4, 383–403.
- GRUBBS, F. 1969. Procedures for detecting outlying observations in samples. *Technometrics* 11, 1, 1–21.
- GUHA, S., RASTOGI, R., AND SHIM, K. 2000. ROCK: A robust clustering algorithm for categorical attributes. *Inform. Syst.* 25, 5, 345–366.
- GUNTER, S., SCHRAUDOLPH, N. N., AND VISHWANATHAN, S. V. N. 2007. Fast iterative kernel principal component analysis. *J. Mach. Learn. Res.* 8, 1893–1918.
- GUTTORSSON, S. E., MARKS R. J. II, EL-SHARKAWI, M. A., AND KERSZENBAUM, I. 1999. Elliptical novelty grouping for online short-turn detection of excited running rotors. *IEEE Trans. Energy Conv.* 14, 1.
- GWADERA, R., ATALLAH, M. J., AND SZPANKOWSKI, W. 2004. Detection of significant sets of episodes in event sequences. In *Proceedings of the 4th IEEE International Conference on Data Mining*. IEEE Computer Society, 3–10.
- GWADERA, R., ATALLAH, M. J., AND SZPANKOWSKI, W. 2005a. Markov models for identification of significant episodes. In *Proceedings of the 5th SIAM International Conference on Data Mining*.
- GWADERA, R., ATALLAH, M. J., AND SZPANKOWSKI, W. 2005b. Reliable detection of episodes in event sequences. *Knowl. Inform. Syst.* 7, 4, 415–437.
- HARRIS, T. 1993. Neural network in machine health monitoring. *Professional Engin.*
- HARTIGAN, J. A. AND WONG, M. A. 1979. A k-means clustering algorithm. *Appl. Stat.* 28, 100–108.
- HAUTAMAKI, V., KARKKAINEN, I., AND FRANTI, P. 2004. Outlier detection using k-nearest neighbour graph. In *Proceedings of the 17th International Conference on Pattern Recognition*. vol. 3. IEEE Computer Society, 430–433.
- HAWKINS, D. 1980. *Identification of Outliers*. Chapman and Hall, London and New York.
- HAWKINS, D. M. 1974. The detection of errors in multivariate data using principal components. *J. Amer. Statist. Assoc.* 69, 346, 340–344.
- HAWKINS, S., HE, H., WILLIAMS, G. J., AND BAXTER, R. A. 2002. Outlier detection using replicator neural networks. In *Proceedings of the 4th International Conference on Data Warehousing and Knowledge Discovery*. Springer-Verlag, 170–180.
- HAZEL, G. G. 2000. Multivariate Gaussian MRF for multi-spectral scene segmentation and anomaly detection. *GeoRS* 38, 3, 1199–1211.
- HE, H., WANG, J., GRACO, W., AND HAWKINS, S. 1997. Application of neural networks to detection of medical fraud. *Expert Syst. Appl.* 13, 4, 329–336.
- HE, Z., DENG, S., AND XU, X. 2002. Outlier detection integrating semantic knowledge. In *Proceedings of the 3rd International Conference on Advances in Web-Age Information Management*. Springer-Verlag, 126–131.
- HE, Z., DENG, S., XU, X., AND HUANG, J. Z. 2006. A fast greedy algorithm for outlier mining. In *Proceedings of the 10th Pacific-Asia Conference on Knowledge and Data Discovery*. 567–576.
- HE, Z., XU, X., AND DENG, S. 2003. Discovering cluster-based local outliers. *Pattern Recog. Lett.* 24, 9–10, 1641–1650.
- HE, Z., XU, X., AND DENG, S. 2005. An optimization model for outlier detection in categorical data. In *Proceedings of the International Conference on Intelligent Computing*. Lecture Notes in Computer Science, vol. 3644. Springer.
- HE, Z., XU, X., HUANG, J. Z., AND DENG, S. 2004a. *A Frequent Pattern Discovery Method for Outlier Detection*. Springer, 726–732.
- HE, Z., XU, X., HUANG, J. Z., AND DENG, S. 2004b. *Mining Class Outliers: Concepts, Algorithms and Applications*. Springer, 588–589.
- HELLER, K. A., SVORE, K. M., KEROMYTIS, A. D., AND STOLFO, S. J. 2003. One class support vector machines for detecting anomalous windows registry accesses. In *Proceedings of the Workshop on Data Mining for Computer Security*.

- HELMAN, P. AND BHANGOO, J. 1997. A statistically-based system for prioritizing information exploration under uncertainty. In *Proceedings of the IEEE International Conference on Systems, Man, and Cybernetics*. vol. 27. IEEE, 449–466.
- HELMER, G., WONG, J., HONAVAR, V., AND MILLER, L. 1998. Intelligent agents for intrusion detection. In *Proceedings of the IEEE Information Technology Conference*. 121–124.
- HICKINBOTHAM, S. J. AND AUSTIN, J. 2000a. Novelty detection in airframe strain data. In *Proceedings of the 15th International Conference on Pattern Recognition*. Vol. 2. 536–539.
- HICKINBOTHAM, S. J. AND AUSTIN, J. 2000b. Novelty detection in airframe strain data. In *Proceedings of the IEEE-INNS-ENNS International Joint Conference on Neural Networks*. vol. 6. 24–27.
- HO, L. L., MACEY, C. J., AND HILLER, R. 1999. A distributed and reliable platform for adaptive anomaly detection in IP networks. In *Proceedings of the 10th IFIP/IEEE International Workshop on Distributed Systems: Operations and Management*. Springer-Verlag, 33–46.
- HO, T. V. AND ROUAT, J. 1997. A novelty detector using a network of integrate and fire neurons. *Lecture Notes in Computer Science*, vol. 1327. 103–108.
- HO, T. V. AND ROUAT, J. 1998. Novelty detection based on relaxation time of a network of integrate-and-fire neurons. In *Proceedings of the 2nd IEEE World Congress on Computational Intelligence*. 1524–1529.
- HODGE, V. AND AUSTIN, J. 2004. A survey of outlier detection methodologies. *Artif. Intel. Rev.* 22, 2, 85–126.
- HOFMEYR, S. A., FORREST, S., AND SOMAYAJI, A. 1998. Intrusion detection using sequences of system calls. *J. Comput. Secur.* 6, 3, 151–180.
- HOLLIER, G. AND AUSTIN, J. 2002. Novelty detection for strain-gauge degradation using maximally correlated components. In *Proceedings of the European Symposium on Artificial Neural Networks*. 257–262–539.
- HOLLMER, J. AND TRESP, V. 1999. Call-based fraud detection in mobile communication networks using a hierarchical regime-switching model. In *Proceedings of the Conference on Advances in Neural Information Processing Systems II*. MIT Press, 889–895.
- HORN, P. S., FENG, L., LI, Y., AND PESCE, A. J. 2001. Effect of outliers and nonhealthy individuals on reference interval estimation. *Clinical Chem.* 47, 12, 2137–2145.
- HU, W., LIAO, Y., AND VEMURI, V. R. 2003. Robust anomaly detection using support vector machines. In *Proceedings of the International Conference on Machine Learning*. Morgan Kaufmann Publishers Inc., 282–289.
- HUBER, P. 1974. *Robust Statistics*. Wiley, New York.
- HUBER, P. J. 1985. Projection pursuit (with discussions). *Ann. Stat.* 13, 2, 435–475.
- IDE, T. AND KASHIMA, H. 2004. Eigenspace-based anomaly detection in computer systems. In *Proceedings of the 10th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 440–449.
- IDÉ, T., PAPADIMITRIOU, S., AND VLACHOS, M. 2007. Computing correlation anomaly scores using stochastic nearest neighbors. In *Proceedings of the International Conference Data Mining*. 523–528.
- IHLER, A., HUTCHINS, J., AND SMYTH, P. 2006. Adaptive event detection with time-varying Poisson processes. In *Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 207–216.
- ILGUN, K., KEMMERER, R. A., AND PORRAS, P. A. 1995. State transition analysis: A rule-based intrusion detection approach. *IEEE Trans. Softw. Eng.* 21, 3, 181–199.
- JAGADISH, H. V., KOUDAS, N., AND MUTHUKRISHNAN, S. 1999. Mining deviants in a time series database. In *Proceedings of the 25th International Conference on Very Large Data Bases*. Morgan Kaufmann Publishers Inc., 102–113.
- JAGOTA, A. 1991. Novelty detection on a very large number of memories stored in a hopfield-style network. In *Proceedings of the International Joint Conference on Neural Networks*. vol. 2. 905.
- JAIN, A. K. AND DUBES, R. C. 1988. *Algorithms for Clustering Data*. Prentice-Hall, Inc.
- JAKUBEK, S. AND STRASSER, T. 2002. Fault-diagnosis using neural networks with ellipsoidal basis functions. In *Proceedings of the American Control Conference*. vol. 5. 3846–3851.
- JANAKIRAM, D., REDDY, V., AND KUMAR, A. 2006. Outlier detection in wireless sensor networks using Bayesian belief networks. In *Proceedings of the 1st International Conference on Communication System Software and Middleware*. 1–6.
- JAPKOWICZ, N., MYERS, C., AND GLUCK, M. A. 1995. A novelty detection approach to classification. In *Proceedings of the International Joint Conference on Artificial Intelligence*. 518–523.
- JAVITZ, H. S. AND VALDES, A. 1991. The SRI IDES statistical anomaly detector. In *Proceedings of the IEEE Symposium on Research in Security and Privacy*. IEEE Computer Society.

- JIANG, M. F., TSENG, S. S., AND SU, C. M. 2001. Two-phase clustering process for outliers detection. *Patt. Recog. Lett.* 22, 6-7, 691–700.
- JIN, W., TUNG, A. K. H., AND HAN, J. 2001. Mining top-n local outliers in large databases. In *Proceedings of the 7th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 293–298.
- JOACHIMS, T. 2006. Training linear SVMs in linear time. In *Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM, 217–226.
- JOLLIFFE, I. T. 2002. *Principal Component Analysis*, 2nd Ed. Springer.
- JOSHI, M. V., AGARWAL, R. C., AND KUMAR, V. 2001. Mining needle in a haystack: classifying rare classes via two-phase rule induction. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*. ACM Press, 91–102.
- JOSHI, M. V., AGARWAL, R. C., AND KUMAR, V. 2002. Predicting rare classes: can boosting make any weak learner strong? In *Proceedings of the 8th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM, 297–306.
- KADOTA, K., TOMINAGA, D., AKIYAMA, Y., AND TAKAHASHI, K. 2003. Detecting outlying samples in micro-array data: A critical assessment of the effect of outliers on sample classification. *Chem-Bio Informatics* 3, 1, 30–45.
- KARYPIS, G. AND KUMAR, V. 1998. Multi-level k-way partitioning scheme for irregular graphs. *J. Paral. Distrib. Comput.* 48, 1, 96–129.
- KEARNS, M. J. 1990. *Computational Complexity of Machine Learning*. MIT Press.
- KEOGH, E., LIN, J., LEE, S.-H., AND HERLE, H. V. 2006. Finding the most unusual time series subsequence: Algorithms and applications. *Knowl. Inform. Syst.* 11, 1, 1–27.
- KEOGH, E., LONARDI, S., AND CHI' CHIU, B. Y. 2002. Finding surprising patterns in a time series database in linear time and space. In *Proceedings of the 8th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 550–556.
- KEOGH, E., LONARDI, S., AND RATANAMAHATANA, C. A. 2004. Towards parameter-free data mining. In *Proceedings of the 10th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 206–215.
- KEOGH, E. AND SMYTH, P. 1997. A probabilistic approach to fast pattern matching in time series databases. In *Proceedings of the 3rd International Conference on Knowledge Discovery and Data Mining*, D. Heckerman, H. Mannila, D. Pregibon, and R. Uthurusamy, Eds. AAAI Press, 24–30.
- KING, S., KING, D., P. ANUZIS, K. A., TARASSENKO, L., HAYTON, P., AND UTETE, S. 2002. The use of novelty detection techniques for monitoring high-integrity plant. In *Proceedings of the International Conference on Control Applications*. vol. 1., 221–226.
- KITAGAWA, G. 1979. On the use of AIC for the detection of outliers. *Technometrics* 21, 2, 193–199.
- KNORR, E. M. AND NG, R. T. 1997. A unified approach for mining outliers. In *Proceedings of the Conference of the Centre for Advanced Studies on Collaborative Research*. IBM Press, 11.
- KNORR, E. M. AND NG, R. T. 1998. Algorithms for mining distance-based outliers in large datasets. In *Proceedings of the 24th International Conference on Very Large Data Bases*. Morgan Kaufmann Publishers Inc., 392–403.
- KNORR, E. M. AND NG, R. T. 1999. Finding intensional knowledge of distance-based outliers. *VLDB J.* 211–222.
- KNORR, E. M., NG, R. T., AND TUCAKOV, V. 2000. Distance-based outliers: Algorithms and applications. *VLDB J.* 8, 3-4, 237–253.
- KO, H. AND JACYNA, G. 2000. Dynamical behavior of autoassociative memory performing novelty filtering. In *IEEE Trans. Neural Netw.* Vol. 11. 1152–1161.
- KOHONEN, T., Ed. 1997. *Self-Organizing Maps*. Springer-Verlag.
- KOJIMA, K. AND ITO, K. 1999. Autonomous learning of novel patterns by utilizing chaotic dynamics. In *Proceedings of the IEEE International Conference on Systems, Man, and Cybernetics*. Vol. 1. IEEE, 284–289.
- KOSORESOW, A. P. AND HOFMEYER, S. A. 1997. Intrusion detection via system call traces. *IEEE Softw.* 14, 5, 35–42.
- KOU, Y., LU, C.-T., AND CHEN, D. 2006. Spatial weighted outlier detection. In *Proceedings of the SIAM Conference on Data Mining*.
- KRUEGEL, C., MUTZ, D., ROBERTSON, W., AND VALEUR, F. 2003. Bayesian event classification for intrusion detection. In *Proceedings of the 19th Annual Computer Security Applications Conference*. IEEE Computer Society, 14.

- KRUEGEL, C., TOTH, T., AND KIRDA, E. 2002. Service specific anomaly detection for network intrusion detection. In *Proceedings of the ACM symposium on Applied Computing*. ACM Press, 201–208.
- KRUEGEL, C. AND VIGNA, G. 2003. Anomaly detection of Web-based attacks. In *Proceedings of the 10th ACM Conference on Computer and Communications Security*. ACM Press, 251–261.
- KUMAR, V. 2005. Parallel and distributed computing for cybersecurity. *IEEE Distrib. Syst. Online* 6, 10.
- LABIB, K. AND VEMURI, R. 2002. NSOM: A real-time network-based intrusion detection using self-organizing maps. *Netw. Security*.
- LAKHINA, A., CROVELLA, M., AND DIOT, C. 2005. Mining anomalies using traffic feature distributions. In *Proceedings of the Conference on Applications, Technologies, Architectures, and Protocols for Computer Communications*. ACM Press, 217–228.
- LANE, T. AND BRODLEY, C. E. 1997a. An application of machine learning to anomaly detection. In *Proceedings of the Conference on 20th NIST-NCSC National Information Systems Security Conference*. 366–380.
- LANE, T. AND BRODLEY, C. E. 1997b. Sequence matching and learning in anomaly detection for computer security. In *Proceedings of the Conference on AI Approaches to Fraud Detection and Risk Management*, Fawcett, Haimowitz, Provost, and Stolfo, Eds. AAAI Press, 43–49.
- LANE, T. AND BRODLEY, C. E. 1999. Temporal sequence learning and data reduction for anomaly detection. *ACM Trans. Inform. Syst. Secur.* 2, 3, 295–331.
- LAUER, M. 2001. A mixture approach to novelty detection using training data with outliers. In *Proceedings of the 12th European Conference on Machine Learning*. Springer-Verlag, 300–311.
- LAURIKKALA, J., JUHOLA, M., AND KENTALA, E. 2000. Informal identification of outliers in medical data. In *Proceedings of the 5th International Workshop on Intelligent Data Analysis in Medicine and Pharmacology*. 20–24.
- LAZAREVIC, A., ERTÖZ, L., KUMAR, V., ÖZGÜR, A., AND SRIVASTAVA, J. 2003. A comparative study of anomaly detection schemes in network intrusion detection. In *Proceedings of the SIAM International Conference on Data Mining*. (SIAM).
- LEE, W. AND STOLFO, S. 1998. Data mining approaches for intrusion detection. In *Proceedings of the 7th USENIX Security Symposium*.
- LEE, W., STOLFO, S., AND CHAN, P. 1997. Learning patterns from UNIX process execution traces for intrusion detection. In *Proceedings of the AAAI Workshop on AI Methods in Fraud and Risk Management*.
- LEE, W., STOLFO, S. J., AND MOK, K. W. 2000. Adaptive intrusion detection: A data mining approach. *Artif. Intell. Rev.* 14, 6, 533–567.
- LEE, W. AND XIANG, D. 2001. Information-theoretic measures for anomaly detection. In *Proceedings of the IEEE Symposium on Security and Privacy*. IEEE Computer Society, 130.
- LI, M. AND VITANYI, P. M. B. 1993. *An Introduction to Kolmogorov Complexity and Its Applications*. Springer-Verlag.
- LI, Y., PONT, M. J., AND JONES, N. B. 2002. Improving the performance of radial basis function classifiers in condition monitoring and fault diagnosis applications where unknown faults may occur. *Patt. Recog. Lett.* 23, 5, 569–577.
- LIN, J., KEOGH, E., FU, A., AND HERLE, H. V. 2005. Approximations to magic: Finding unusual medical time series. In *Proceedings of the 18th IEEE Symposium on Computer-Based Medical Systems*. IEEE Computer Society, 329–334.
- LIN, S. AND BROWN, D. E. 2003. An outlier-based data association method for linking criminal incidents. In *Proceedings of the 3rd SIAM Data Mining Conference*.
- LIU, J. P. AND WENG, C. S. 1991. Detection of outlying data in bioavailability/bioequivalence studies. *Stat. Med.* 10, 9, 1375–89.
- LU, C.-T., CHEN, D., AND KOU, Y. 2003. Algorithms for spatial outlier detection. In *Proceedings of the 3rd International Conference on Data Mining*. 597–600.
- MA, J. AND PERKINS, S. 2003a. Online novelty detection on temporal sequences. In *Proceedings of the 9th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 613–618.
- MA, J. AND PERKINS, S. 2003b. Time-series novelty detection using one-class support vector machines. In *Proceedings of the International Joint Conference on Neural Networks*. Vol. 3. 1741–1745.
- MACDONALD, J. W. AND GHOSH, D. 2007. Copa-cancer outlier profile analysis. *Bioinformatics* 22, 23, 2950–2951.
- MAHONEY, M. V. AND CHAN, P. K. 2002. Learning nonstationary models of normal network traffic for detecting novel attacks. In *Proceedings of the 8th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 376–385.

- MAHONEY, M. V. AND CHAN, P. K. 2003. Learning rules for anomaly detection of hostile network traffic. In *Proceedings of the 3rd IEEE International Conference on Data Mining*. IEEE Computer Society, 601.
- MAHONEY, M. V., CHAN, P. K., AND ARSHAD, M. H. 2003. A machine learning approach to anomaly detection. Tech. rep. CS-2003-06, Department of Computer Science, Florida Institute of Technology Melbourne.
- MANEVITZ, L. M. AND YOUSEF, M. 2000. Learning from positive data for document classification using neural networks. In *Proceedings of the 2nd Bar-Ilan Workshop on Knowledge Discovery and Learning*.
- MANEVITZ, L. M. AND YOUSEF, M. 2002. One-class SVMs for document classification. *J. Mach. Learn. Res.* 2, 139–154.
- MANIKOPOULOS, C. AND PAPAVASSILIOU, S. 2002. Network intrusion and fault detection: A statistical anomaly approach. *IEEE Comm. Mag.* 40.
- MANSON, G. 2002. Identifying damage sensitive, environment insensitive features for damage detection. In *Proceedings of IES Conference*.
- MANSON, G., PIERCE, G., AND WORDEN, K. 2001. On the long-term stability of normal conditions for damage detection in a composite panel. In *Proceedings of the 4th International Conference on Damage Assessment of Structures*. Cardiff, UK.
- MANSON, G., PIERCE, S. G., WORDEN, K., MONNIER, T., GUY, P., AND ATHERTON, K. 2000. Long-term stability of normal condition data for novelty detection. In *Proceedings of the Conference on Smart Structures and Integrated Systems*. 323–334.
- MARCEAU, C. 2000. Characterizing the behavior of a program using multiple-length n-grams. In *Proceedings of the Workshop on New Security Paradigms*. ACM Press, 101–110.
- MARCHETTE, D. 1999. A statistical method for profiling network traffic. In *Proceedings of the 1st USENIX Workshop on Intrusion Detection and Network Monitoring*. 119–128.
- MARKOU, M. AND SINGH, S. 2003a. Novelty detection: A review-part 1: Statistical approaches. *Sig. Proc.* 83, 12, 2481–2497.
- MARKOU, M. AND SINGH, S. 2003b. Novelty detection: A review-part 2: Neural network based approaches. *Sig. Proc.* 83, 12, 2499–2521.
- MARSLAND, S., NEHMZOW, U., AND SHAPIRO, J. 1999. A model of habituation applied to mobile robots. In *Proceedings of Towards Intelligent Mobile Robots Conference*. Department of Computer Science, Manchester University, Technical rep. UMCS-99-3-1.
- MARSLAND, S., NEHMZOW, U., AND SHAPIRO, J. 2000a. Novelty detection for robot neotaxis. In *Proceedings of the 2nd International Symposium on Neural Computation*. 554–559.
- MARSLAND, S., NEHMZOW, U., AND SHAPIRO, J. 2000b. A real-time novelty detector for a mobile robot. In *Proceedings of the EUREL Conference on Advanced Robotics Systems*.
- MARTINELLI, G. AND PERFETTI, R. 1994. Generalized cellular neural network for novelty detection. *IEEE Trans. Circ. Syst. I: Fundamental Theory Application* 41, 2, 187–190.
- MARTINEZ, D. 1998. Neural tree density estimation for novelty detection. *IEEE Trans. Neural Netw.* 9, 2, 330–338.
- MCCALLUM, A., NIGAM, K., AND UNGAR, L. H. 2000. Efficient clustering of high-dimensional data sets with application to reference matching. In *Proceedings of the 6th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 169–178.
- MCNEIL, A. 1999. Extreme value theory for risk managers. In *Internal Modelling and CAD II*, 93–113.
- MINGMING, N. Y. 2000. Probabilistic networks with undirected links for anomaly detection. In *Proceedings of the IEEE Systems, Man, and Cybernetics Information Assurance and Security Workshop*. 175–179.
- MOTULSKY, H. 1995. *Intuitive Biostatistics: Choosing a Statistical Test*. Oxford University Press, Chapter 37.
- MOYA, M., KOCH, M., AND HOSTETLER, L. 1993. One-class classifier networks for target recognition applications. In *Proceedings of the World Congress on Neural Networks, International Neural Network Society*. 797–801.
- MURRAY, A. F. 2001. Novelty detection using products of simple experts: A potential architecture for embedded systems. *Neural Netw.* 14, 9, 1257–1264.
- NAIRAC, A., CORBETT-CLARK, T., RIPLEY, R., TOWNSEND, N., AND TARASSENKO, L. 1997. Choosing an appropriate model for novelty detection. In *Proceedings of the 5th IEEE International Conference on Artificial Neural Networks*. 227–232.
- NAIRAC, A., TOWNSEND, N., CARR, R., KING, S., COWLEY, P., AND TARASSENKO, L. 1999. A system for the analysis of jet engine vibration data. *Integ. Comput.-Aided Eng.* 6, 1, 53–56.
- NG, R. T. AND HAN, J. 1994. Efficient and effective clustering methods for spatial data mining. In *Proceedings*

- of the 20th International Conference on Very Large Data Bases. Morgan Kaufmann Publishers Inc., 144–155.
- NOBLE, C. C. AND COOK, D. J. 2003. Graph-based anomaly detection. In *Proceedings of the 9th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 631–636.
- ODIN, T. AND ADDISON, D. 2000. Novelty detection using neural network technology. In *Proceedings of the COMADEN Conference*.
- OTTEY, M., PARTHASARATHY, S., GHOTING, A., LI, G., NARRAVULA, S., AND PANDA, D. 2003. Towards NIC-based intrusion detection. In *Proceedings of the 9th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 723–728.
- OTTEY, M. E., GHOTING, A., AND PARTHASARATHY, S. 2006. Fast distributed outlier detection in mixed-attribute data sets. *Data Min. Knowl. Disc.* 12, 2-3, 203–228.
- PALSHIKAR, G. K. 2005. Distance-based outliers in sequences. *Lecture Notes in Computer Science*, vol. 3816, 547–552.
- PAPADIMITRIOU, S., KITAGAWA, H., GIBBONS, P. B., AND FALOUTSOS, C. 2002. Loci: Fast outlier detection using the local correlation integral. Tech. rep. IRP-TR-02-09, Intel Research Laboratory.
- PARRA, L., DECO, G., AND MIESBACH, S. 1996. Statistical independence and novelty detection with information preserving nonlinear maps. *Neural Comput.* 8, 2, 260–269.
- PARZEN, E. 1962. On the estimation of a probability density function and mode. *Annals Math. Stat.* 33, 1065–1076.
- PATCHA, A. AND PARK, J.-M. 2007. An overview of anomaly detection techniques: Existing solutions and latest technological trends. *Comput. Netw.* 51, 12, 3448–3470.
- PETSCHKE, T., MARCANTONIO, A., DARKEN, C., HANSON, S., KUHN, G., AND SANTOSO, I. 1996. A neural network autoassociator for induction motor failure prediction. In *Proceedings of the Conference on Advances in Neural Information Processing*. vol. 8, 924–930.
- PHOHA, V. V. 2002. *The Springer Internet Security Dictionary*. Springer-Verlag.
- PHUA, C., ALAHAKOON, D., AND LEE, V. 2004. Minority report in fraud detection: Classification of skewed data. *SIGKDD Explorer Newsletter* 6, 1, 50–59.
- PHUONG, T. V., HUNG, L. X., CHO, S. J., LEE, Y., AND LEE, S. 2006. An anomaly detection algorithm for detecting attacks in wireless sensor networks. *Intel. Secur. Inform.* 3975, 735–736.
- PICKANDS, J. 1975. Statistical inference using extreme order statistics. *Annals Stat.* 3, 1, 119–131.
- PIRES, A. AND SANTOS-PEREIRA, C. 2005. Using clustering and robust estimators to detect outliers in multivariate data. In *Proceedings of the International Conference on Robust Statistics*.
- PLATT, J. 2000. Probabilistic Outputs for Support Vector Machines and Comparison to Regularized Likelihood Methods. In *Advances in Large Margin Classifiers*, A. Smola, P. Bartlett, B. Schoelkopf, and D. Schuurmans, Eds. MIT Press, 61–74.
- POKRAJAC, D., LAZAREVIC, A., AND LATECKI, L. J. 2007. Incremental local outlier detection for data streams. In *Proceedings of the IEEE Symposium on Computational Intelligence and Data Mining*.
- PORRAS, P. A. AND NEUMANN, P. G. 1997. EMERALD: Event monitoring enabling responses to anomalous live disturbances. In *Proceedings of the 20th NIST-NCSC National Information Systems Security Conference*. 353–365.
- PORTNOY, L., ESKIN, E., AND STOLFO, S. 2001. Intrusion detection with unlabeled data using clustering. In *Proceedings of the ACM Workshop on Data Mining Applied to Security*.
- PROTOPAPAS, P., GIAMMARCO, J. M., FACCIOLO, L., STRUBLE, M. F., DAVE, R., AND ALCOCK, C. 2006. Finding outlier light curves in catalogues of periodic variable stars. *Monthly Notices Royal Astronomical Soc.* 369, 2, 677–696.
- QIN, M. AND HWANG, K. 2004. Frequent episode rules for Internet anomaly detection. In *Proceedings of the 3rd IEEE International Symposium on Network Computing and Applications*. IEEE Computer Society.
- RAMADAS, M., OSTERMANN, S., AND TJADEN, B. C. 2003. Detecting anomalous network traffic with self-organizing maps. In *Proceedings of the Conference on Recent Advances in Intrusion Detection*. 36–54.
- RAMASWAMY, S., RASTOGI, R., AND SHIM, K. 2000. Efficient algorithms for mining outliers from large data sets. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*. ACM Press, 427–438.
- RATSCH, G., MIKA, S., SCHOLKOPF, B., AND MULLER, K.-R. 2002. Constructing boosting algorithms from SVMs: An application to one-class classification. *IEEE Trans. Patt. Anal. Mach. Intel.* 24, 9, 1184–1199.
- ROBERTS, S. 1999. Novelty detection using extreme value statistics. In *Proceedings of the IEEE Vision, Image and Signal Processing Conference* Vol. 146. 124–129.

- ROBERTS, S. 2002. Extreme value statistics for novelty detection in biomedical signal processing. In *Proceedings of the 1st International Conference on Advances in Medical Signal and Information Processing*. 166–172.
- ROBERTS, S. AND TARASSENKO, L. 1994. A probabilistic resource allocating network for novelty detection. *Neural Comput.* 6, 2, 270–284.
- ROSNER, B. 1983. Percentage points for a generalized ESD many-outlier procedure. *Technometrics* 25, 2, 165–172.
- ROTH, V. 2004. Outlier detection with one-class kernel Fisher discriminants. In *Proceedings of the Conference on Advances in Neural Information Processing Systems (NIPS)*.
- ROTH, V. 2006. Kernel fisher discriminants for outlier detection. *Neural Comput.* 18, 4, 942–960.
- ROUSSEEUW, P. J. AND LEROY, A. M. 1987. *Robust Regression and Outlier Detection*. John Wiley & Sons, Inc.
- ROUSSOPOULOS, N., KELLEY, S., AND VINCENT, F. 1995. Nearest neighbor queries. In *Proceedings of the ACM-SIGMOD International Conference on Management of Data*.
- RUOTOLO, R. AND SURACE, C. 1997. A statistical approach to damage detection through vibration monitoring. In *Proceedings of the 5th Pan-American Congress of Applied Mechanics*.
- SALVADOR, S. AND CHAN, P. 2003. Learning states and rules for time-series anomaly detection. Tech. rep. CS-2003-05, Department of Computer Science, Florida Institute of Technology Melbourne.
- SARAWAGI, S., AGRAWAL, R., AND MEGIDDO, N. 1998. Discovery-driven exploration of OLAP data cubes. In *Proceedings of the 6th International Conference on Extending Database Technology*. Springer-Verlag, 168–182.
- SARGOR, C. 1998. Statistical anomaly detection for link-state routing protocols. In *Proceedings of the 6th International Conference on Network Protocols*. IEEE Computer Society, 62.
- SAUNDERS, R. AND GERO, J. 2000. The importance of being emergent. In *Proceedings of the Conference on Artificial Intelligence in Design*.
- SCARTH, G., MCINTYRE, M., WOWK, B., AND SOMORJAI, R. 1995. Detection of novelty in functional images using fuzzy clustering. In *Proceedings of the 3rd Meeting of the International Society for Magnetic Resonance in Medicine*. 238.
- SCHÖLKOPF, B., PLATT, J. C., SHAW-TAYLOR, J. C., SMOLA, A. J., AND WILLIAMSON, R. C. 2001. Estimating the support of a high-dimensional distribution. *Neural Comput.* 13, 7, 1443–1471.
- SCOTT, S. L. 2001. Detecting network intrusion using a Markov modulated nonhomogeneous Poisson Process. *Journal of the American Statistical Association*.
- SEBYALA, A. A., OLUKEMI, T., AND SACKS, L. 2002. Active platform security through intrusion detection using naive Bayesian network for anomaly detection. In *Proceedings of the London Communications Symposium*.
- SEQUEIRA, K. AND ZAKI, M. 2002. Admit: Anomaly-based data mining for intrusions. In *Proceedings of the 8th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 386–395.
- SHEIKHOLESAMI, G., CHATTERJEE, S., AND ZHANG, A. 1998. Wavecluster: A multi-resolution clustering approach for very large spatial databases. In *Proceedings of the 24th International Conference on Very Large Databases*. Morgan Kaufmann Publishers Inc., 428–439.
- SHEKHAR, S., LU, C.-T., AND ZHANG, P. 2001. Detecting graph-based spatial outliers: Algorithms and applications (a summary of results). In *Proceedings of the 7th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 371–376.
- SHEWHART, W. A. 1931. *Economic Control of Quality of Manufactured Product*. D. Van Nostrand Company.
- SHYU, M.-L., CHEN, S.-C., SARINNAKORN, K., AND CHANG, L. 2003. A novel anomaly detection scheme-based on principal component classifier. In *Proceedings of the 3rd IEEE International Conference on Data Mining*. 353–365.
- SIATERLIS, C. AND MAGLARIS, B. 2004. Towards multi-sensor data fusion for dos detection. In *Proceedings of the ACM Symposium on Applied Computing*. ACM Press, 439–446.
- SINGH, S. AND MARKOU, M. 2004. An approach to novelty detection applied to the classification of image regions. *IEEE Trans. Knowl. Data Eng.* 16, 4, 396–407.
- SMITH, R., BIVENS, A., EMBRECHTS, M., PALAGIRI, C., AND SZYMANSKI, B. 2002. Clustering approaches for anomaly-based intrusion detection. In *Proceedings of the Intelligent Engineering Systems through Artificial Neural Networks*. ASME Press, 579–584.
- SMYTH, P. 1994. Markov monitoring with unknown states. *IEEE J. Select. Areas Comm.* (Special Issue on Intelligent Signal Processing for Communications) 12, 9, 1600–1612.
- SNYDER, D. 2001. Online intrusion detection using sequences of system calls. M.S. thesis, Department of Computer Science, Florida State University.

- SOHN, H., WORDEN, K., AND FARRAR, C. 2001. Novelty detection under changing environmental conditions. In *Proceedings of the 8th Annual SPIE International Symposium on Smart Structures and Materials*.
- SOLBERG, H. E. AND LAHTI, A. 2005. Detection of outliers in reference distributions: Performance of Horn's algorithm. *Clinical Chem.* 51, 12, 2326–2332.
- SONG, Q., HU, W., AND XIE, W. 2002. Robust support vector machine with bullet hole image classification. *IEEE Trans. Syst. Man Cyber.—Part C: Applications and Reviews* 32, 4.
- SONG, S., SHIN, D., AND YOON, E. 2001. Analysis of novelty detection properties of auto-associators. In *Proceedings of the Conference on Condition Monitoring and Diagnostic Engineering Management*. 577–584.
- SONG, X., WU, M., JERMAINE, C., AND RANKA, S. 2007. Conditional anomaly detection. *IEEE Trans. Knowl. Data Eng.* 19, 5, 631–645.
- SOULE, A., SALAMATIAN, K., AND TAFT, N. 2005. Combining filtering and statistical methods for anomaly detection. In *Proceedings of the 5th ACM SIGCOMM Conference on Internet Measurement*. ACM, 1–14.
- SPENCE, C., PARRA, L., AND SAJDA, P. 2001. Detection, synthesis and compression in mammographic image analysis with a hierarchical image probability model. In *Proceedings of the IEEE Workshop on Mathematical Methods in Biomedical Image Analysis*. IEEE Computer Society, 3.
- SRIVASTAVA, A. 2006. Enabling the discovery of recurring anomalies in aerospace problem reports using high-dimensional clustering techniques. In *Proceedings of the IEEE Aerospace Conference*, 17–34.
- SRIVASTAVA, A. AND ZANE-ULMAN, B. 2005. Discovering recurring anomalies in text reports regarding complex space systems. In *Proceedings of the IEEE Aerospace Conference*, 3853–3862.
- STEFANO, C., SANSONE, C., AND VENTO, M. 2000. To reject or not to reject: that is the question: An answer in the case of neural classifiers. *IEEE Trans. Syst. Manag. Cyber.* 30, 1, 84–94.
- STEFANSKY, W. 1972. Rejecting outliers in factorial designs. *Technometrics* 14, 2, 469–479.
- STEINWART, I., HUSH, D., AND SCOVEL, C. 2005. A classification framework for anomaly detection. *J. Mach. Learn. Res.* 6, 211–232.
- STREIFEL, R., MAK, R., AND EL-SHARKAWI, M. 1996. Detection of shorted-turns in the field of turbine-generator rotors using novelty detectors—development and field tests. *IEEE Trans. Energy Conv.* 11, 2, 312–317.
- SUBRAMANIAM, S., PALPANAS, T., PAPADOPOULOS, D., KALOGERAKI, V., AND GUNOPULOS, D. 2006. Online outlier detection in sensor data using non-parametric models. In *Proceedings of the 32nd International Conference on Very Large Data Bases (VLDB)*. VLDB Endowment, 187–198.
- SUN, H., BAO, Y., ZHAO, F., YU, G., AND WANG, D. 2004. CD-trees: An efficient index structure for outlier detection. In *Proceedings of the 5th International Conference on Web-Age Information Management (WAIM)*. 600–609.
- SUN, J., QU, H., CHAKRABARTI, D., AND FALOUTSOS, C. 2005. Neighborhood formation and anomaly detection in bipartite graphs. In *Proceedings of the 5th IEEE International Conference on Data Mining*. IEEE Computer Society, 418–425.
- SUN, J., XIE, Y., ZHANG, H., AND FALOUTSOS, C. 2007. Less is more: Compact matrix representation of large sparse graphs. In *Proceedings of the 7th SIAM International Conference on Data Mining*.
- SUN, P. AND CHAWLA, S. 2004. On local spatial outliers. In *Proceedings of the 4th IEEE International Conference on Data Mining*. 209–216.
- SUN, P. AND CHAWLA, S. 2006. SLOM: A new measure for local spatial outliers. *Knowl. Inform. Syst.* 9, 4, 412–429.
- SUN, P., CHAWLA, S., AND ARUNASALAM, B. 2006. Mining for outliers in sequential databases. In *Proceedings of the SIAM International Conference on Data Mining*.
- SURACE, C. AND WORDEN, K. 1998. A novelty detection method to diagnose damage in structures: An application to an offshore platform. In *Proceedings of the 8th International Conference of Off-Shore and Polar Engineering*. vol. 4. Colorado, 64–70.
- SURACE, C., WORDEN, K., AND TOMLINSON, G. 1997. A novelty detection approach to diagnose damage in a cracked beam. In *Proceedings of the SPIE*. vol. 3089. 947–953.
- SUZUKI, E., WATANABE, T., YOKOI, H., AND TAKABAYASHI, K. 2003. Detecting interesting exceptions from medical test data with visual summarization. In *Proceedings of the 3rd IEEE International Conference on Data Mining*. 315–322.
- SYKACEK, P. 1997. Equivalent error bars for neural network classifiers trained by Bayesian inference. In *Proceedings of the European Symposium on Artificial Neural Networks*. 121–126.
- TAN, P.-N., STEINBACH, M., AND KUMAR, V. 2005. *Introduction to Data Mining*. Addison-Wesley.
- TANDON, G. AND CHAN, P. 2007. Weighting versus pruning in rule validation for detecting network and host

- anomalies. In *Proceedings of the 13th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press.
- TANG, J., CHEN, Z., CHEE FU, A. W., AND W.CHEUNG, D. 2002. Enhancing effectiveness of outlier detections for low density patterns. In *Proceedings of the Pacific-Asia Conference on Knowledge Discovery and Data Mining*. 535–548.
- TANIGUCHI, M., HAFT, M., HOLLMN, J., AND TRESP, V. 1998. Fraud detection in communications networks using neural and probabilistic methods. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing*. vol. 2. IEEE Computer Society, 1241–1244.
- TAO, Y., XIAO, X., AND ZHOU, S. 2006. Mining distance-based outliers from large databases in any metric space. In *Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 394–403.
- TARASSENKO, L. 1995. Novelty detection for the identification of masses in mammograms. In *Proceedings of the 4th IEEE International Conference on Artificial Neural Networks*. vol. 4. 442–447.
- TAX, D. AND DUIN, R. 1999a. Data domain description using support vectors. In *Proceedings of the European Symposium on Artificial Neural Networks*, M. Verleysen, Ed., 251–256.
- TAX, D. AND DUIN, R. 1999b. Support vector data description. *Patt. Recog. Lett.* 20, 11–13, 1191–1199.
- TAX, D. M. J. 2001. One-class classification; concept-learning in the absence of counter-examples. Ph.D. thesis, Delft University of Technology.
- TENG, H., CHEN, K., AND LU, S. 1990. Adaptive real-time anomaly detection using inductively generated sequential patterns. In *Proceedings of the IEEE Computer Society Symposium on Research in Security and Privacy*. IEEE Computer Society Press, 278–284.
- THEILER, J. AND CAI, D. M. 2003. Resampling approach for anomaly detection in multispectral images. In *Proceedings of the SPIE*. vol. 5093, 230–240.
- THOMPSON, B., II, R. M., CHOI, J., EL-SHARKAWI, M., HUANG, M., AND BUNJE, C. 2002. Implicit learning in auto-encoder novelty assessment. In *Proceedings of the International Joint Conference on Neural Networks*. 2878–2883.
- THOTTAN, M. AND JI, C. 2003. Anomaly detection in IP networks. *IEEE Trans. Sig. Proc.* 51, 8, 2191–2204.
- TIBSHIRANI, R. AND HASTIE, T. 2007. Outlier sums for differential gene expression analysis. *Biostatistics* 8, 1, 2–8.
- TOMLINS, S. A., RHODES, D. R., PERNER, S., DHANASEKARAN, S. M., MEHRA, R., SUN, X. W., VARAMBALLY, S., CAO, X., TCHINDA, J., KUEFER, R., LEE, C., MONTIE, J. E., SHAH, R., PIENTA, K. J., RUBIN, M., AND CHINNAIYAN, A. M. 2005. Recurrent fusion of tmprss2 and ets transcription factor genes in prostate cancer. *Science* 310, 5748, 603–611.
- TORR, P. AND MURRAY, D. 1993. Outlier detection and motion segmentation. In *Proceedings of the SPIE. Sensor Fusion VI*, S. Schenker, Ed. vol. 2059. 432–443.
- TSAY, R. S., PEA, D., AND PANKRATZ, A. E. 2000. Outliers in multi-variate time series. *Biometrika* 87, 4, 789–804.
- VAIDYA, J. AND CLIFTON, C. 2004. Privacy-preserving outlier detection. In *Proceedings of the 4th IEEE International Conference on Data Mining*. 233–240.
- VALDES, A. AND SKINNER, K. 2000. Adaptive, model-based monitoring for cyber attack detection. In *Proceedings of the 3rd International Workshop on Recent Advances in Intrusion Detection*. Springer-Verlag, 80–92.
- VAPNIK, V. N. 1995. *The Nature of Statistical Learning Theory*. Springer-Verlag.
- VASCONCELOS, G., FAIRHURST, M., AND BISSET, D. 1994. Recognizing novelty in classification tasks. In *Proceedings of the Neural Information Processing Systems Workshop on Novelty Detection and Adaptive Systems Monitoring*.
- VASCONCELOS, G. C., FAIRHURST, M. C., AND BISSET, D. L. 1995. Investigating feed-forward neural networks with respect to the rejection of spurious patterns. *Patt. Recog. Lett.* 16, 2, 207–212.
- VILALTA, R. AND MA, S. 2002. Predicting rare events in temporal domains. In *Proceedings of the IEEE International Conference on Data Mining*. IEEE Computer Society, 474.
- VINUEZA, A. AND GRUDIC, G. 2004. Unsupervised outlier detection and semi-supervised learning. Tech. rep. CU-CS-976-04, University of Colorado at Boulder.
- WEI, L., QIAN, W., ZHOU, A., AND JIN, W. 2003. Hot: Hypergraph-based outlier test for categorical data. In *Proceedings of the 7th Pacific-Asia Conference on Knowledge and Data Discovery*. 399–410.
- WEIGEND, A. S., MANGEAS, M., AND SRIVASTAVA, A. N. 1995. Nonlinear gated experts for time-series: Discovering regimes and avoiding overfitting. *Int. J. Neural Syst.* 6, 4, 373–399.

- WEISS, G. M. AND HIRSH, H. 1998. Learning to predict rare events in event sequences. In *Proceedings of the 4th International Conference on Knowledge Discovery and Data Mining*, R. Agrawal, P. Stolorz, and G. Piatetsky-Shapiro, Eds. AAAI Press, 359–363.
- WHITEHEAD, B. AND HOYT, W. 1993. A function approximation approach to anomaly detection in propulsion system test data. In *Proceedings of the 29th AIAA/SAE/ASME/ASEE Joint Propulsion Conference*. IEEE Computer Society.
- WILLIAMS, G., BAXTER, R., HE, H., HAWKINS, S., AND GU, L. 2002. A comparative study of RNN for outlier detection in data mining. In *Proceedings of the IEEE International Conference on Data Mining*. IEEE Computer Society, 709.
- WONG, W.-K., MOORE, A., COOPER, G., AND WAGNER, M. 2002. Rule-based anomaly pattern detection for detecting disease outbreaks. In *Proceedings of the 18th National Conference on Artificial Intelligence*. MIT Press. <http://www.cs.cmu.edu/simawm/antiterror>.
- WONG, W.-K., MOORE, A., COOPER, G., AND WAGNER, M. 2003. Bayesian network anomaly pattern detection for disease outbreaks. In *Proceedings of the 20th International Conference on Machine Learning*. AAAI Press, 808–815.
- WORDEN, K. 1997. Structural fault detection using a novelty measure. *J. Sound Vibr.* 201, 1, 85–101.
- WU, M. AND JERMAINE, C. 2006. Outlier detection by sampling with accuracy guarantees. In *Proceedings of the 12th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM, 767–772.
- WU, N. AND ZHANG, J. 2003. Factor analysis based anomaly detection. In *Proceedings of the IEEE Workshop on Information Assurance*. United States Military Academy.
- YAIRI, T., KATO, Y., AND HORI, K. 2001. Fault detection by mining association rules from housekeeping data. In *Proceedings of the International Symposium on Artificial Intelligence, Robotics and Automation in Space*.
- YAMANISHI, K. AND ICHI TAKEUCHI, J. 2001. Discovering outlier filtering rules from unlabeled data: Combining a supervised learner with an unsupervised learner. In *Proceedings of the 7th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM Press, 389–394.
- YAMANISHI, K., TAKEUCHI, J.-I., WILLIAMS, G., AND MILNE, P. 2004. Online unsupervised outlier detection using finite mixtures with discounting learning algorithms. *Data Min. Knowl. Disc.* 8, 275–300.
- YE, N. AND CHEN, Q. 2001. An anomaly detection technique based on a chi-square statistic for detecting intrusions into information systems. *Quality Reliability Engin. Int.* 17, 105–112.
- YI, B.-K., SIDIROPOULOS, N., JOHNSON, T., JAGADISH, H. V., FALOUTSOS, C., AND BILIRIS, A. 2000. Online data mining for co-evolving time sequences. In *Proceedings of the 16th International Conference on Data Engineering*. IEEE Computer Society, 13.
- YPM, A. AND DUIN, R. 1998. Novelty detection using self-organizing maps. In *Progress in Connectionist Based Information Systems*. vol. 2. Springer, 1322–1325.
- YU, D., SHEIKHOLESAMI, G., AND ZHANG, A. 2002. Findout: Finding outliers in very large datasets. *Knowl. Inform. Syst.* 4, 4, 387–412.
- YU, J. X., QIAN, W., LU, H., AND ZHOU, A. 2006. Finding centric local outliers in categorical/numerical spaces. *Knowl. Inform. Syst.* 9, 3, 309–338.
- ZEEVI, A. J., MEIR, R., AND ADLER, R. 1997. Time series prediction using mixtures of experts. In *Advances in Neural Information Processing*. vol. 9. MIT Press.
- ZHANG, J. AND WANG, H. 2006. Detecting outlying subspaces for high-dimensional data: The new task, algorithms, and performance. *Knowl. Inform. Syst.* 10, 3, 333–355.
- ZHANG, K., SHI S., GAO, H., AND LI, J. 2007. Unsupervised outlier detection in sensor networks using aggregation tree. In *Advanced Data Mining and Applications 4632*, 158–169.
- ZHANG, Z., LI, J., MANIKOPOULOS, C., JORGENSEN, J., AND UCLES, J. 2001. Hide: A hierarchical network intrusion detection system using statistical preprocessing and neural network classification. In *Proceedings of the IEEE Workshop on Information Assurance and Security*. West Point, 85–90.
- ZIMMERMANN, J. AND MOHAY, G. 2006. Distributed intrusion detection in clusters based on non-interference. In *Proceedings of the Australasian Workshops on Grid Computing and E-Research (ACSW Frontiers)*. Australian Computer Society, Inc., 89–95.

Received November 2007; revised March 2008; accepted May 2008