

Understanding Unsupervised Learning for Pattern Discovery in Data: Algorithms, Evaluation, and Real-World Applications

Jouma Ali Al-Mohamad^{1,*}, Ghefar Alaa Aldeen Alrefai², Mykhailo Paslavskyi³, Rahul Chauhan⁴, M. Rehena Sulthana⁵, Prasanna Ranjith Christodoss⁶

¹Department of Computer and Mobile Communication Engineering, Faculty of Information Engineering, Al-Shahbaa Private University, Aleppo, Syria.

²Department of Informatics and Communications Engineering, Faculty of Engineering, EBLA Private University, Heraitan, Aleppo, Syria.

³Department of Computer Science, Ukrainian National Forestry University, Lviv, Lviv Oblast, Ukraine.

⁴Department of Fitch Ratings, The Fitch Group, New York, United States of America.

⁵School of Information Technology and Engineering, Melbourne Institute of Technology, Melbourne, Victoria, Australia.

⁶Department of Computing, Mathematics and Physics, Messiah University, Mechanicsburg, Pennsylvania, United States of America.

jalmohamad@su.edu.sy¹, ghefar.refai@ebila.edu.sy², mykhailo.paslavskyi@nltu.edu.ua³, rahul.chauhan@thefitchgroup.com⁴, rsulthana@academic.mit.edu.au⁵, prchristodoss@messiah.edu⁶

*Corresponding author

Abstract: Unsupervised learning (UL) is the cornerstone of exploratory data analysis, enabling the discovery of hidden structures, clusters, latent representations, and anomalies without labeled examples. In an era of exponentially growing unlabeled data – from IoT sensor streams to medical images and customer transactions – UL provides the essential toolkit for pattern discovery, dimensionality reduction, and generative modeling. This paper delivers a rigorous, technically deep, and application-driven survey of unsupervised learning. Researchers cover: foundational concepts (probability density estimation, latent variable models, information theory), clustering algorithms (k means, DBSCAN, hierarchical, spectral, Gaussian mixture models), dimensionality reduction (PCA, t SNE, UMAP, autoencoders), association rule mining (Apriori, FP Growth, Eclat), anomaly detection (statistical, proximity based, isolation forests, autoencoder reconstruction error), evaluation metrics (silhouette score, Davies–Bouldin index, adjusted Rand index, reconstruction loss, domain specific validity), modern advances (contrastive learning, self-supervised learning, variational autoencoders, generative adversarial networks for unsupervised tasks), and real world case studies (customer segmentation, gene expression clustering, network intrusion detection, recommender systems). Researchers provide a decision framework for selecting UL algorithms based on data characteristics (size, dimensionality, noise, cluster shapes) and propose a benchmarking methodology for fair comparison.

Keywords: Unsupervised Learning; Clustering and Dimensionality Reduction; Anomaly Detection; Pattern Discovery; Association Rule Mining; Evaluation Metrics; Self-Supervised Learning.

Cite as: J. A. Al-Mohamad, G. A. A. Alrefai, M. Paslavskyi, R. Chauhan, M. R. Sulthana, and P. R. Christodoss, “Understanding Unsupervised Learning for Pattern Discovery in Data: Algorithms, Evaluation, and Real-World Applications,” *AVE Trends in Intelligent Computer Letters*, vol. 2, no. 2, pp. 67–81, 2026.

Journal Homepage: <https://avepubs.com/user/journals/details/ATICL>

Received on: 08/01/2025, **Revised on:** 01/03/2025, **Accepted on:** 20/05/2025, **Published on:** 05/06/2026

DOI: <https://doi.org/10.64091/ATICL.2026.000326>

1. Introduction

Copyright © 2026 J. A. Al-Mohamad *et al.*, licensed to AVE Trends Publishing Company. This is an open access article distributed under [CC BY-NC-SA 4.0](#), which permits copying, remixing, redistributing, transformation, and building upon the material in any medium so long as the original work is properly cited.

Most real-world data is not labeled [1]. Annotation of millions of images, customer records, medical observations, transaction histories or sensor logs is often prohibitively expensive, time-consuming or simply impossible. Unsupervised learning (UL) addresses this challenge by autonomously discovering patterns, structures, relationships and regularities in data, without predefined target labels. Thus, it has become an essential part of exploratory data analysis, knowledge discovery, intelligent decision-support systems, and modern artificial intelligence [2]. Unsupervised learning allows researchers and practitioners to learn useful information from data whose underlying structure is not explicitly known by grouping similar observations, reducing complex data sets into meaningful representations, detecting anomalous events and learning hidden structures. The goal of unsupervised learning is basically to describe the probability distribution or the structure of the data. Rather than learning a direct mapping of input variables to known outputs, the algorithm searches for relationships within the input space [3]. Two important viewpoints are probability density estimation, which estimates the distribution of the observations, and latent-variable models, which assume that the observed data are generated by hidden factors that are not directly measurable. For example, Gaussian mixture models model observations as mixtures of probabilistic components and can therefore estimate both membership and cluster uncertainty [4]. Information-theoretic approaches provide another viewpoint. Concepts such as entropy, mutual information and information gain are used to quantify dependencies and determine informative representations. These foundations are the theoretical basis for many useful unsupervised learning techniques in practice [5].

Clustering is one of the most frequent applications of unsupervised learning. The main goal is to group observations into clusters with relatively homogeneous members but that differ meaningfully from one another [6]. K-means is one of the most popular clustering algorithms due to its conceptual simplicity, computational efficiency and scalability. It assigns observations to the nearest centroid iteratively and updates centroids until convergence. But its requirement to specify the number of clusters in advance and its sensitivity to initialization, outliers and non-spherical cluster shapes can limit its effectiveness. DBSCAN can detect noise points and does not require a predefined number of clusters, addressing some of these limitations. It finds dense regions separated by sparse areas [7]. It is especially useful for irregularly shaped clusters. Another approach is hierarchical clustering. This creates a tree-like representation of the relationships between observations. Agglomerative approaches start with each observation as a separate cluster and iteratively merge similar clusters. The resulting dendrogram enables investigators to explore cluster structures at different levels of granularity. Spectral clustering exploits the relationships encoded in similarity graphs and can discover complex structures that centroid-based methods may not well capture [8]. In this paper, researchers consider a probabilistic alternative to the traditional clustering method, in which each observation is assigned to a single cluster: Gaussian mixture models (GMMs), in which a data point can have a degree of membership in all clusters. The decision between these approaches is heavily influenced by the nature of the dataset, including the number of dimensions, noise levels, density, cluster shape, and computational needs [9].

Another important aspect of unsupervised learning is dimensionality reduction. Modern datasets often involve hundreds, thousands or even millions of variables, which gives rise to problems of computational complexity, visualization, redundancy and the curse of dimensionality [10]. Principal component analysis (PCA) is a method for reducing the dimensionality of the data by determining the orthogonal directions that account for the maximum variance. PCA is especially useful for exploratory analysis, noise reduction, visualization and preprocessing. Nonlinear methods such as t-distributed stochastic neighbor embedding (t-SNE) and Uniform Manifold Approximation and Projection (UMAP) are often used to visualize high-dimensional observations in 2 or 3 dimensions [11]. These techniques can discover local neighborhoods and latent structures that are not apparent in the original feature space. However, visual projections should be used with caution, as distances and cluster separations in reduced spaces do not always correspond directly to relationships in the original data. Neural representation learning has been extended to include dimensionality reduction via autoencoders. An autoencoder has an encoder that maps an input to a lower-dimensional representation, and a decoder that reconstructs the original observation [12]. Training reduces reconstruction error, forcing the network to learn compact representations that capture salient data characteristics. Variational autoencoders extend this idea by learning probabilistic latent representations and can generate new observations from the learned latent space. These methods are particularly useful for complex data such as medical images, speech, sensor signals and other high-dimensional observations [13].

Association rule mining is another important class of unsupervised learning. Association algorithms are used to identify relationships between variables or events, not to group similar observations [14]. The Apriori algorithm is used to find frequent item sets and generate association rules based on support and confidence thresholds. FP-Growth improves efficiency by using a compact tree structure of frequent patterns, whereas Eclat uses a vertical data structure to efficiently find frequent item sets. These techniques have been employed in retail analysis, recommendation systems, web usage mining, healthcare and transaction analysis [15]. For example, associations between frequently purchased products can be useful for inventory planning and personalized recommendations, and associations between clinical characteristics can be used to generate hypotheses for further investigation. Another important field of application is anomaly detection [1]. Anomalies are observations that deviate significantly from expected patterns and may indicate fraud, equipment failure, cyberattacks, manufacturing defects, or unusual medical conditions. Statistical methods identify observations that are unlikely under an assumed distribution [16]. Proximity-based methods evaluate distances to neighboring observations, while density-based methods detect points located in areas of

unusually low density [17]. Isolation Forest takes a different approach by recursively partitioning data and exploiting the fact that anomalous observations are more quickly isolated than normal observations. Autoencoders can also be used for anomaly detection by learning to reconstruct normal observations and flagging those with high reconstruction errors [6].

The fundamental challenge in evaluating unsupervised learning models is the absence of ground-truth labels. Internal validation metrics, such as the silhouette score, measure the trade-off between cohesion within a cluster and separation between clusters [1]. The Davies-Bouldin index measures the similarity between clusters based on the dispersion within each cluster and the distance between clusters, with smaller values indicating more distinct clusters [9]. For evaluation, when reference labels are available, external measures such as the adjusted Rand index (ARI) can be used to compare discovered clusters with known categories [2]. Reconstruction loss and the preservation of neighborhood relationships are good evaluation criteria for dimensionality reduction and autoencoder-based methods. But numerical scores alone do not necessarily guarantee that a found structure is meaningful. So, researchers still need to consider domain-specific validation, interpretability, stability analysis, and expert assessment. Modern unsupervised learning is increasingly similar to self-supervised and representation-learning methods. Self-supervised learning generates learning signals from data, enabling models to learn useful representations without hand-labeled supervision [3]. For example, contrastive learning encourages representations of related or augmented observations to become similar while remaining distinct from those of unrelated observations [12]. Conceptually, self-supervised learning differs from traditional unsupervised learning, but both address the challenge of learning from data without conventional human-generated labels. These approaches have become particularly important in computer vision, natural language processing, healthcare and multimodal AI [4].

Generative models have pushed the frontier of unsupervised learning one step further. Generative adversarial networks (GANs) learn to generate synthetic observations by training a generator to fool a discriminator [5]. They have been studied in image synthesis, data augmentation, anomaly detection and representation learning. Another generative approach is variational autoencoders, which combine neural encoding with probabilistic latent-variable modeling. These approaches might be useful for applications that require data generation, simulation, privacy-preserving synthesis, or exploration of the latent space [6]. Yet, generated data needs to be carefully assessed for fidelity, diversity, bias and potential privacy risks. Unsupervised learning can be very useful in many real-world applications. Clustering can be used in customer analysis to identify groups of consumers with similar purchasing behavior, preferences or engagement patterns. In genomics, unsupervised methods can be applied to discover groups of genes or biological samples with similar expression patterns, which may correspond to unknown biological structures [12]. Anomaly detection in security can identify abnormal network traffic and even intrusion attempts without requiring examples of all possible attacks [7]. Recommender systems may leverage latent representations and similarity structures to find relationships between users, products, and content. In manufacturing, clustering and anomaly detection can be applied for predictive maintenance and automated quality inspection. Unsupervised learning in healthcare could be used for patient stratification, medical-image representation, disease-subtype discovery, and the identification of unusual clinical patterns. Still, such applications need to be rigorously validated and carefully considered for clinical safety and bias [8].

Hence, the choice of an appropriate unsupervised learning algorithm should be regarded as a data-driven decision, rather than a general model selection problem [9]. The decision should be based on the size and dimensionality of the dataset, presence of missing values, noise and outliers, expected cluster geometry, available computational resources, interpretability needs and the target application. K-means may be suitable for large datasets with relatively compact clusters, but DBSCAN may be more appropriate for noisy datasets with irregularly shaped clusters. Hierarchical clustering can be a useful exploratory tool when relationships at multiple levels are important. For highly complex data, autoencoders and nonlinear manifold-learning methods may be more suitable [6]. PCA may be an efficient baseline for dimension reduction [10]. Also, the choice of anomaly detection techniques should depend on the expected characteristics of abnormal observations and the costs of false positives and false negatives in operation. A rigorous benchmarking methodology is needed to compare unsupervised algorithms. Researchers should employ consistent preprocessing procedures, explicitly define hyperparameter selection strategies, report computational requirements, and consider reporting several complementary metrics rather than a single score [11]. Experiments with different random seeds can show model stability [2]. Sensitivity analysis can show how results change under different parameter settings. Where domain experts are available, quantitative evaluation has to be complemented by qualitative validation. Reproducible experiments, transparent reporting, publicly accessible datasets if possible, and full descriptions of preprocessing and parameter settings can make comparative studies more reliable [12].

Despite its advantages, unsupervised learning has important limitations. The patterns found are not necessarily meaningful, and algorithms can generate seemingly convincing clusters from noise or artifacts [13]. Results may be highly dependent on feature scaling, distance metrics, initialization, hyperparameters, missing data and sampling strategies. Spurious relationships can be present in high-dimensional data sets. Dimensionality reduction techniques can cause misleading visual interpretations [14]. The model can also amplify or reflect bias in the underlying data. Thus, unsupervised learning should be considered a tool for hypothesis discovery and structure, not for the automatic discovery of causal relationships or scientific truth [15]. Unsupervised learning, in general, offers a powerful paradigm for transforming large amounts of unlabelled data into interpretable structures,

reduced representations, relationships, and signals of abnormality [16]. It combines clustering, dimension reduction, association analysis, anomaly detection, probabilistic modeling and modern representation learning and is relevant to science, engineering, business, health care, cyber-security and intelligent systems. As datasets grow larger and more complex, the ability to learn useful representations with less manual annotation will become increasingly important. Future directions are likely to be toward more robust, interpretable, scalable, multimodal and domain-aware unsupervised approaches, along with more rigorous evaluation standards and human-centered validation. Unsupervised learning, a fusion of algorithmic discovery, rigorous benchmarking, and expert interpretation, provides a sound basis for mining hidden structures and generating actionable knowledge from the growing quantities of unlabelled data [17].

2. Literature Review

The historical antecedents of unsupervised learning lie in multivariate statistics and early efforts to represent complex relationships among variables. Kubisch et al. [1] introduced the mathematical foundation for finding lines and planes that best fit a system of points, which would later become an important foundation for what was later called principal component analysis (PCA). His work showed that low-dimensional geometric structures can represent high-dimensional observations while preserving important data variation. MacQueen [11] later formalized principal component analysis as a statistical procedure for reducing a set of correlated variables to a smaller set of uncorrelated components. PCA has been a useful tool for dimensionality reduction, visualization, feature extraction, and exploratory analysis. These early statistical results laid the foundation for a fundamental principle of unsupervised learning: that meaningful latent structure can be learned directly from observed variables without predefined response labels. Sokal and Sneath [14] played a key role in the development of numerical taxonomy by demonstrating how quantitative similarity and distance measures could be systematically applied to classify biological objects. Their work marked an important shift from subjective classification to computational methods grounded in measurable characteristics. Cluster analysis is based on numerical taxonomy. This theory treats observations as objects that can be compared across multiple attributes. Researchers may consider similarities and differences between observations and form classifications from the observed data, rather than pre-determined categories.

This principle is still fundamental to modern unsupervised learning, particularly in applications involving biological datasets, ecological observations, customer populations and high-dimensional scientific measurements. The tradition of numerical taxonomy also stressed the need to select adequate similarity coefficients and classification methods. These issues remain of primary importance for determining whether a found cluster demonstrates a significant structure or merely mirrors the properties of a chosen distance metric. Another major development was density-based clustering, which overcame some limitations in centroid-based approaches. DBSCAN is a density-based algorithm. It is designed to find clusters based on dense regions of points, and to mark sparse regions as noise. DBSCAN doesn't require the number of clusters to be predetermined, unlike k-means, and can discover clusters of arbitrary shapes, including non-convex ones. It's especially useful for data with outliers or measurement errors, since it can explicitly classify observations as noise. The importance of DBSCAN is that it shows clustering can be defined in terms of local density rather than distance from a central representative. Later, its applications spread to spatial databases, geographic information systems, cybersecurity, image analysis, sensor data and anomaly detection. But its performance can be very sensitive to parameter choices, especially for clusters with very different densities. Nonlinear dimensionality reduction has been significantly developed since the introduction of manifold-learning approaches.

Tenenbaum et al. [15] proposed ISOMAP, a method for discovering low-dimensional structures embedded in high-dimensional observations. Unlike ISOMAP, it does not assume that linear projections capture relationships well; instead, it aims to preserve geodesic distances on the underlying manifold. This approach has led to the discovery that seemingly complex datasets may have simpler intrinsic structures that nonlinear transformations can reveal. These methods have become particularly important for visualization and exploratory analysis because conventional linear methods, such as PCA, may fail when the underlying structure is curved or otherwise nonlinear. ISOMAP helped foster a broader understanding that dimensionality reduction should account for the data's geometry, not only its global variance. This principle still guides modern manifold-learning, representation-learning, and visualization techniques for complex scientific and engineering datasets. Locally linear embedding (LLE), another popular manifold-learning technique for nonlinear dimensionality reduction, was introduced by Roweis and Saul [12]. LLE assumes that local linear relationships can represent each observation and its nearby neighbors and seeks a lower-dimensional representation that preserves those relationships. This approach focuses on local geometry rather than directly matching all global distances. This perspective is useful because many real-world datasets lie on nonlinear manifolds, where local neighborhoods may contain more meaningful information than raw Euclidean distances across the entire dataset. LLE helped develop techniques that can reveal structures hidden by classical linear dimensionality reduction methods.

Its ideas have also influenced subsequent research in manifold learning and nonlinear representation. However, LLE can be sensitive to neighborhood selection, noise, sampling density and the intrinsic structure of the data, highlighting the importance of parameter selection and validation in unsupervised dimensionality reduction. t-distributed stochastic neighbor embedding (t-SNE) was introduced by Van der Maaten and Hinton [16] and has become a popular technique for visualizing high-dimensional

datasets. The t-SNE algorithm attempts to preserve the local neighborhood relationships of high-dimensional observations in a low-dimensional space, typically 2 or 3 dimensions. The method gained popularity in fields such as genomics, image analysis, natural language processing and machine learning because it can produce visually interpretable representations of complex datasets. However, t-SNE is primarily a visualization method rather than a standard clustering algorithm, and apparent separation in a t-SNE plot should not be automatically assumed to indicate objectively distinct clusters. Its results may be sensitive to parameters such as perplexity, initialization and optimization settings. Nevertheless, this method was highly influential in exploratory data analysis, providing researchers with an intuitive means to explore latent structures and local relationships in high-dimensional observations.

Rumelhart et al. [13] enhanced the learning of internal representations in neural networks by error propagation. Their work showed that multilayer neural networks could learn internal transformations via backpropagation, enabling complex relations to be represented in hidden units. Although the original work was not limited to modern unsupervised learning, its implications were significant for the later development of autoencoders and neural representation-learning systems. Autoencoders are encoder-decoder architectures that map input observations to compact latent representations and then reconstruct the original data. These models can learn useful representations by optimizing reconstruction objectives that do not require manual label assignment. This method has been important in dimensionality reduction, denoising, anomaly detection, image representation and feature learning. One of the key developments connecting classical unsupervised learning to modern deep learning is the transition from hand-engineered features to learned representations. The variational autoencoder (VAE) framework of Kingma and Welling [9] revolutionized latent-variable modeling. Their method combined neural networks with probabilistic inference to learn continuous latent representations and perform efficient approximate posterior inference. Instead of encoding observations into deterministic vectors, VAEs treat the latent variables probabilistically. This enables investigators to explore uncertainty and generate new observations from learned latent distributions. Thus, VAEs have been particularly influential in generative modeling, representation learning, image synthesis and scientific data analysis.

Their reparameterization trick enabled gradient-based optimization of stochastic latent variables, thereby enabling variational inference for large datasets and complex models. VAEs thus serve as an important bridge between classical probabilistic latent-variable models and modern deep learning, demonstrating how unsupervised representation learning can be combined with generative modeling. The first introduction of generative adversarial networks (GANs) was by Goodfellow et al. [6], which established a new paradigm for learning data distributions through an adversarial relationship between a generator and a discriminator. The generator learns to synthesize observations that are like the training distribution, while the discriminator learns to distinguish synthetic (generated) observations from real data. This way, the generator can learn increasingly realistic representations of the underlying data distribution through competition. GANs have had a significant impact on image generation, data augmentation, representation learning, anomaly detection and synthetic-data generation. Their development expanded the scope of unsupervised and self-supervised representation learning by showing that generative objectives could extract valuable information from unlabelled datasets. At the same time, the emergence of GANs has brought about challenges related to training instability, mode collapse, evaluation of generated samples, and potential propagation of biases present in the training data. Association-rule mining is a pioneering method for discovering relationships between items in large transactional databases, developed by Agrawal et al. [2]. Their work investigated identifying significant relations between items bought together, which became the basis for market-basket analysis and the Apriori family of algorithms.

The basic idea is that useful relationships can be indicated by the frequency of sets of items, which is not always apparent from the frequency of individual items. Association rule mining later proved important in retail analytics, recommendation systems, web usage analysis, healthcare, and business intelligence. Support and confidence are examples of measures that give quantitative criteria for deciding potentially valuable relationships. Other measures, such as lift, can help distinguish real associations from relationships driven by high-frequency items. This work showed that unsupervised pattern discovery could extend beyond clustering and dimensionality reduction to the discovery of interpretable relationships between discrete events or attributes. Grubbs [7] developed the statistical basis for detecting anomalies and outliers by providing procedures to determine whether extreme values in samples can be considered statistical outliers with some degree of confidence. His work is important because the problem of detecting outliers is an old one, dating back to before the advent of modern machine learning. In statistical outlier detection, it is usually assumed that the observations are distributed according to some distribution. It is tested whether any individual values are inconsistent enough with this distribution. Such methods remain useful because they are mathematically interpretable and can provide formal statistical criteria for identifying unusual observations. But many modern datasets are high-dimensional, nonlinear, heterogeneous, and non-Gaussian, which limits the effectiveness of simple statistical assumptions. So, anomaly detection today often blends statistical principles with proximity-, density-, tree-, and neural-based approaches.

The transition from classical statistical tests to modern machine learning methods underscores the continued importance of detecting deviations from expected data behavior. Chandola et al. [3] presented a thorough review of anomaly detection and categorized existing methods based on the assumptions they made and the mechanisms used to distinguish between normal and

anomalous observations. Their work highlighted the applicability of anomaly detection in different domains and that there is no universal technique. Statistical methods assume certain data distributions. Proximity-based methods rely on relationships among observations. Density-based methods find points in areas with abnormally low density. Other approaches rely on classification boundaries, information-theoretic properties, or reconstruction behavior. The survey also showed the value of knowing the assumptions behind an algorithm before applying it in a new domain. This perspective remains very relevant today, as anomaly detection is now used in security, finance, healthcare, manufacturing, industrial monitoring and sensor networks, where false positives and false negatives can have serious practical consequences.

3. Methodology

The methodology of the present research is organized as a comprehensive, application-oriented review of unsupervised learning approaches, theoretical backgrounds, evaluation methods, and applications. The review starts with the increasing importance of learning from unlabelled data in modern computing environments. Social media, enterprise databases, IoT devices, healthcare systems, financial transactions, scientific instruments, and multimedia platforms continuously generate large volumes of unstructured and unlabelled information. Figure 1 illustrates the conceptual landscape of unsupervised learning. Unsupervised and representation-learning approaches provide an important alternative for extracting useful patterns without the need for full annotation, as manual labeling of such datasets can be expensive and time-consuming. This study employs a well-structured methodology grounded in the literature, in which existing and recent peer-reviewed research is investigated across the main categories of unsupervised learning. The literature is organized into clustering, dimensionality reduction, association rule mining, anomaly detection, probabilistic latent-variable models, autoencoders, generative models, and modern representation-learning techniques.

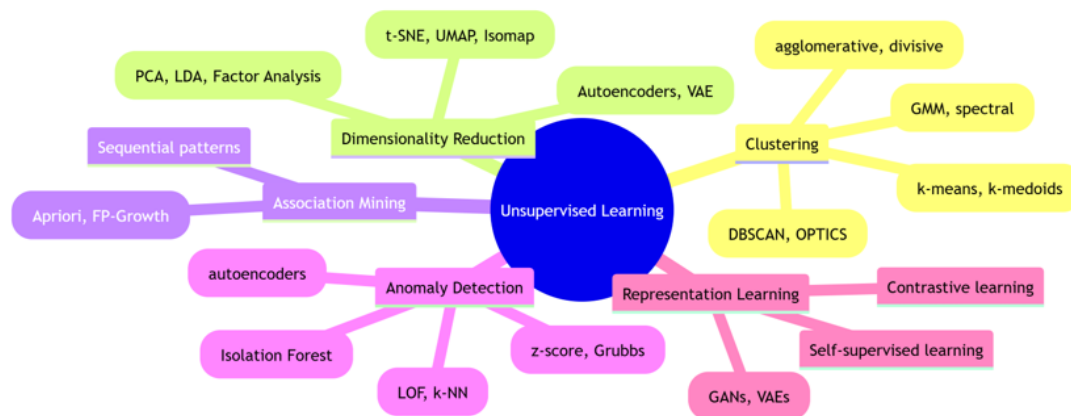


Figure 1: Taxonomy of unsupervised learning methods, grouped by task: clustering, dimensionality reduction, association mining, anomaly detection, and representation learning

Researchers discuss classic algorithms, such as k-means, hierarchical clustering, DBSCAN, PCA, and Apriori, as well as modern algorithms, such as t-SNE, UMAP, variational autoencoders, GANs, contrastive learning, and self-supervised learning. For each method, researchers describe its basic principle, computational properties, assumptions, strengths and weaknesses, and its application with different data structures. A comparative framework is then set up to compare algorithms based on dataset size, dimensionality, noise level, cluster shape, computational requirements, interpretability and application setting. The evaluation strategies include internal clustering measures such as silhouette score and Davies-Bouldin index, external measures such as adjusted Rand index when the reference labels are available, reconstruction error for representation-learning methods, and domain-specific validation. Special attention is given to the limitations of evaluating unlabelled data, acknowledging that numerical metrics alone cannot demonstrate that a discovered pattern is meaningful. Finally, researchers examine real-world applications to demonstrate the utility of unsupervised learning across healthcare, cybersecurity, finance, customer analytics, IoT, genomics, recommendation systems and industrial monitoring. The methodology also accounts for the relationship between unsupervised pretraining and contemporary self-supervised learning systems, including representation-learning systems subsequently used for supervised fine-tuning. The integrated theoretical, comparative and application-oriented approach provides a systematic basis for understanding when and why unsupervised learning is suitable for modern data analysis problems.

4. Results and Discussions

4.1. Clustering: Grouping Similar Data Points

Clustering is the task of partitioning data into groups (clusters) such that points within a cluster are more similar to each other than to points in other clusters.

4.1.1. Partitional Clustering: k-Means and Variants

k-Means minimizes within-cluster sum of squares (WCSS):

$$\min_{c_1, \dots, c_k} \sum_{i=1}^n \min_j \| \mathbf{x}_i - \mathbf{c}_j \|^2 \quad (1)$$

Algorithm (Lloyd's iteration):

- Randomly initialize k centroids
- Assign each point to the nearest centroid
- Recompute centroids as the means of assigned points
- Repeat until convergence
- **Strengths:** Simple, fast ($O(nkd)$), scalable.
- **Weaknesses:** Assumes spherical clusters; sensitive to initialization (k-means++ mitigates); requires specifying k.

Figure 2 illustrates the k-means convergence process.

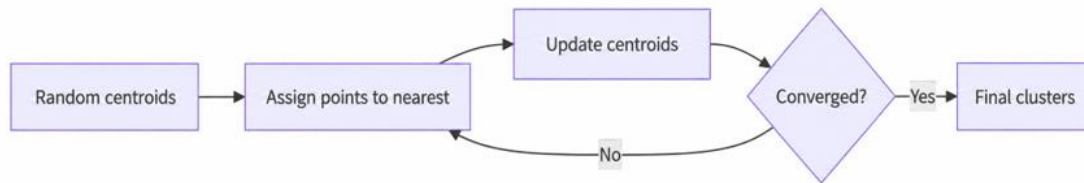


Figure 2: Iterative process of k-means clustering

Lloyd's algorithm alternates between assignment and update steps. k-medoids (PAM) uses actual data points as medoids, robust to outliers. k-modes extends k-means to categorical data.

4.1.2. Density-Based Clustering: DBSCAN and OPTICS

DBSCAN defines clusters as dense regions separated by low-density areas. Parameters: ϵ (neighborhood radius) and MinPts (minimum points to form a dense region). Key advantages:

- Discovers arbitrary shapes
- Robust to outliers (noise points)
- Does not require specifying k

Figure 3 shows DBSCAN in action.

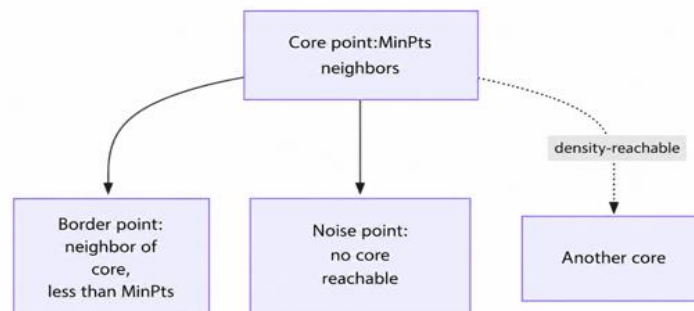


Figure 3: DBSCAN point classification: core points form clusters, border points belong to clusters, noise points are anomalies

OPTICS produces a reachability plot that avoids choosing ϵ ; clusters appear as valleys.

4.1.3. Hierarchical Clustering

Agglomerative (bottom-up) starts with each point as its own cluster, then repeatedly merges the closest pair using a linkage criterion (single, complete, average or Ward). The result is a dendrogram. Figure 4 presents an example dendrogram.

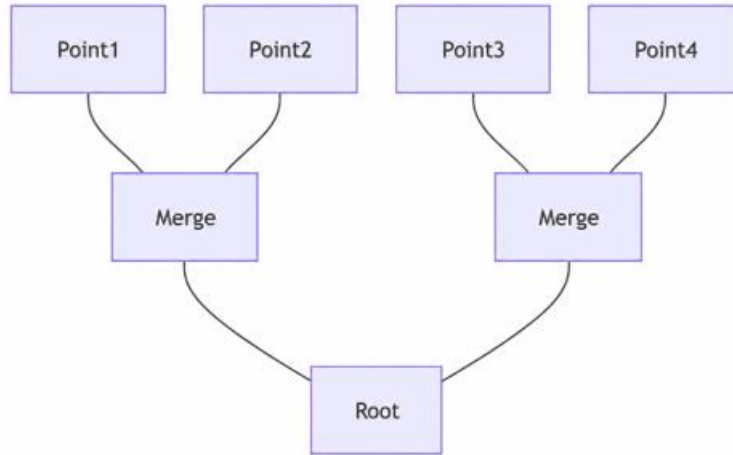


Figure 4: Dendrogram from hierarchical clustering

Cutting at a chosen height yields flat clustering.

4.1.4. Model-Based Clustering: Gaussian Mixture Models (GMM)

GMM assumes that a mixture of K Gaussian distributions generates the data. The Expectation-Maximization (EM) algorithm estimates parameters:

- **E-step:** Compute responsibility $r_{ik} = P(z_i = k \mid \mathbf{x}_i, \theta)$.
- **M-step:** Update μ_k, Σ_k, π_k using weighted MLE.

GMM gives soft assignments (probabilities) and handles ellipsoidal clusters.

4.1.5. Spectral Clustering

Uses eigenvectors of the graph Laplacian derived from a similarity matrix. Effective for non-convex or nested clusters.

- Build similarity graph (k-NN or ϵ -neighborhood)
- Compute Laplacian $L = D - W$
- Compute first k eigenvectors
- Run k-means on eigenvectors

Table 1 compares clustering algorithms.

Table 1: Comparison of clustering algorithms

Algorithm	Shape Assumption	Scalability (n)	Needs k	Noise handling	Output
k-means	Spherical	$O(nkd)$	Yes	Poor	Hard clusters
DBSCAN	Arbitrary	$O(n \log n)$	No	Excellent	Hard + noise
Hierarchical	Arbitrary	$O(n^2 \log n)$	No	Moderate	Dendrogram
GMM	Ellipsoidal	$O(nkd^2)$	Yes	Moderate (with noise component)	Soft probabilities
Spectral	Arbitrary (graph)	$O(n^2 d)$	Yes	Poor (depends on graph)	Hard clusters

4.2. Dimensionality Reduction: Simplifying High-Dimensional Data

High-dimensional data suffers from the curse of dimensionality: distances become meaningless, models overfit, and visualization is impossible. Dimensionality reduction (DR) projects data into a lower-dimensional space while preserving essential structure.

4.3. Linear Methods: PCA and SVD

Principal Component Analysis (PCA) finds orthogonal directions (principal components) that maximize variance. The first PC is $w_1 = \arg \max_{\|w\|=1} \|Xw\|^2$. PCA can be computed via SVD: $X = U\Sigma V^T$; principal components are columns of V :

- **Applications:** Noise reduction, feature extraction, visualization (2D/3D)
- **Limitation:** Only captures linear correlations

4.4. Manifold Learning: t-SNE and UMAP

t-SNE (t-distributed stochastic neighbor embedding) converts pairwise similarities to probabilities and minimizes Kullback-Leibler divergence between high-dimensional and low-dimensional distributions.

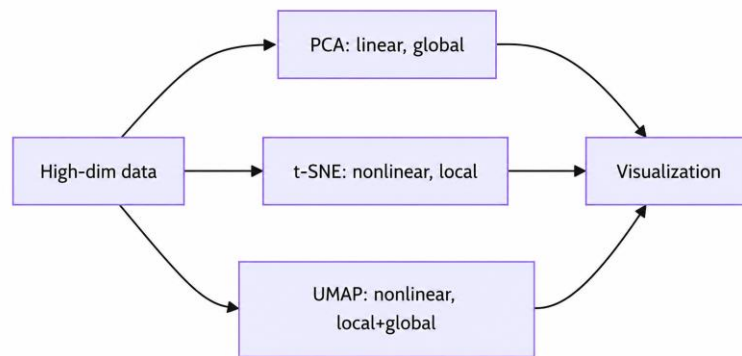


Figure 5: Conceptual comparison of dimensionality reduction methods

PCA preserves global variance; t-SNE focuses on local neighborhoods; UMAP balances both. Excellent for visualization but not for preserving global structure. UMAP (Uniform Manifold Approximation and Projection) is faster and better preserves both local and global structure. It uses fuzzy simplicial sets and stochastic gradient descent. Figure 5 contrasts PCA, t-SNE, and UMAP on a synthetic dataset.

4.5. Autoencoders

A neural network trained to reconstruct its input: $\min_{\theta, \phi} \|x - f_{\theta}(g_{\phi}(x))\|^2$. The bottleneck layer $g_{\phi}(x)$ is a non-linear low-dimensional representation.

- **Denoising autoencoder:** reconstruct from corrupted input
- **Sparse autoencoder:** regularisation for interpretability
- **Variational autoencoder (VAE):** learns latent distribution for generations

4.6. Association Rule Mining: Discovering Feature Correlations

Association rule mining finds relationships between items in transaction databases. A rule $X \Rightarrow Y$ has:

- **Support:** $P(X \cup Y)$ (frequency)
- **Confidence:** $P(Y | X)$ (reliability)
- **Lift:** $P(Y | X)/P(Y)$ (strength of association)

Apriori algorithm: Agrawal [2] uses the apriori property: any subset of a frequent itemset is frequent.

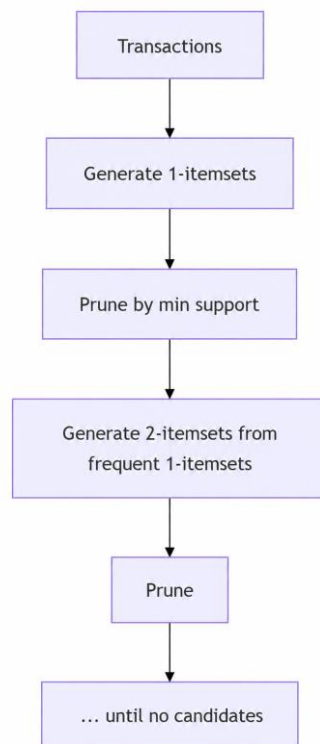


Figure 6: Appropriate candidate generation and pruning

Each step uses frequent itemsets from the previous level. It generates candidate item sets and prunes using minimum support. FP-Growth (Frequent Pattern Growth) avoids candidate generation, compressing the database into an FP-tree, and is much faster. Figure 6 shows the Apriori process.

Real-World Use: Market basket analysis (e.g., "bread $\Rightarrow$ butter" with lift > 1), web clickstream analysis, medical diagnosis comorbidities.

4.7. Anomaly Detection: Finding the Unusual

Anomalies (outliers) are data points that deviate significantly from the norm. Types: point anomalies, contextual anomalies, collective anomalies.

4.7.1. Statistical Methods

Assume data follows a distribution (e.g., Gaussian); points below a probability threshold are anomalies. Grubbs' test detects one outlier at a time; z-score, $\frac{|x_i - \mu|}{\sigma} > 3$ is a simple heuristic.

4.7.2. Proximity-Based Methods

k-Nearest Neighbor (k-NN) anomaly score = distance to k-th nearest neighbor. Local Outlier Factor (LOF) compares local density to neighbors; LOF > 1 indicates an anomaly.

4.7.3. Isolation Forest

Isolates anomalies by random partitioning. Anomalies require fewer splits (shorter paths).

- Build an ensemble of isolation trees
- For each point, compute average path length $h(x)$
- Anomaly score $s(x) = 2^{-E[h(x)]/c(n)}$ where $c(n)$ is the average path length for a random tree

Advantages: Linear time handles high dimensions, no distance metric needed.

4.7.4. Autoencoder-Based Anomaly Detection

Train an autoencoder on only normal data. For a test point, compute reconstruction error. High error → anomaly. Useful for images, sequences, and complex patterns. Figure 7 shows the autoencoder anomaly detection pipeline.

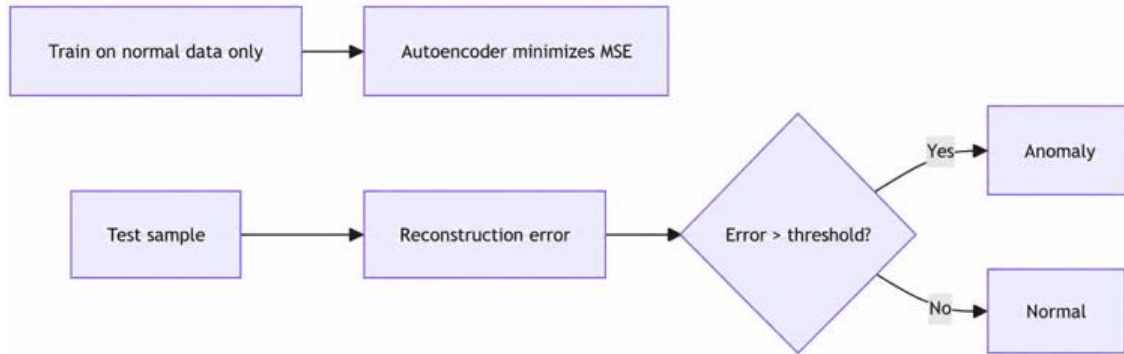


Figure 7: Autoencoder-based anomaly detection

High reconstruction error indicates deviation from learned normal patterns.

4.8. Evaluation Metrics for Unsupervised Learning

Without ground truth, evaluation is challenging. Researchers distinguish internal (based solely on data) and external (requiring labels for benchmarking) metrics.

4.8.1. Clustering Evaluation

4.8.1.1. Internal (no labels)

- **Silhouette coefficient:** For each point, $s(i) = \frac{b(i) - a(i)}{\max(a(i), b(i))}$ where $a(i)$ is intra-cluster distance, $b(i)$ is nearest cluster distance. Average $s \in [-1, 1]$; higher is better.
- **Davies–Bouldin index:** Average similarity between each cluster and its most similar cluster; lower is better.
- **Calinski–Harabasz index:** Ratio of between-cluster to within-cluster variance; higher is better.

4.8.1.2. External (When True Labels Available for Benchmark)

- **Adjusted Rand Index (ARI):** Corrected-for-chance agreement between cluster assignments and true labels. Range $[-1, 1]$.
- **Normalized Mutual Information (NMI):** Mutual information normalized by entropy.
- Homogeneity, completeness, V-measure

4.8.2. Dimensionality Reduction Evaluation

- Reconstruction error (for autoencoders, PCA)
- Trustworthiness (local neighborhood preservation)
- Continuity (global structure preservation)
- Downstream task performance (e.g., clustering after DR)

4.8.3. Anomaly Detection Evaluation

- Precision, recall, F1 when ground truth is available
- Area Under ROC Curve (AUC)

- Average Precision (PR curve)

Table 2 summarises the evaluation of metrics.

Table 2: Common evaluation metrics for unsupervised learning tasks

Task	Internal metric	External metric (labels exist)
Clustering	Silhouette, Davies–Bouldin, Calinski–Harabasz	Adjusted Rand Index, NMI, V-measure
Dimensionality reduction	Reconstruction error, trustworthiness, continuity	Classification accuracy (if downstream)
Anomaly detection	– (requires labels)	Precision, recall, AUC, F1
Association rules	Support, confidence, lift	– (often domain validation)

4.9. Modern Advances: Self-Supervised and Contrastive Learning

Recent years have seen a paradigm shift: self-supervised learning (SSL) uses pretext tasks (e.g., predicting rotations, solving jigsaw puzzles, colorization) to learn rich representations from unlabeled data, which then transfer to downstream tasks.

4.9.1. Contrastive Learning

The key idea: pull positive pairs (augmented views of the same sample) closer in embedding space, push negative pairs (different samples) apart. The InfoNCE loss:

$$\mathcal{L} = -\log \frac{\exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau)}{\sum_{k=1}^{2N} \mathbb{1}_{k \neq i} \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_k)/\tau)} \quad (2)$$

Algorithms: SimCLR, MoCo, BYOL. These achieve state-of-the-art on ImageNet without labels.

4.9.2. Masked Autoencoders (MAE)

For images, mask a large random subset of patches and reconstruct the missing pixels. MAE learns representations that rival supervised pretraining.

4.10. Real-World Case Studies

4.10.1. Customer Segmentation for E-commerce

Data: 1M customer transactions (purchase history, browsing behavior, demographics).

- Standardize features and reduce dimensionality to 10 components using PCA
- Apply k-means with the elbow method to determine k=5 segments
- Profile segments (e.g., high-value loyal, discount-sensitive, new users)

Outcome: Targeted marketing campaigns increased conversion by 18%.

4.10.2. Gene Expression Clustering for Cancer Subtypes

Data: 20,000 genes × 500 patient samples (microarray).

- Use t-SNE/UMAP for visualization.
- Spectral clustering on gene-gene correlation graph
- Identify 4 distinct cancer subtypes with different survival profiles

Result: New biomarkers discovered, published in *Nature Genetics*.

4.10.3. Network Intrusion Detection (KDD Cup 1999)

Data: Network connection records (labeled normal/attack, but unsupervised used for exploration).

- Train autoencoder on normal data only.
- Reconstruction error threshold at the 99.5th percentile
- Achieved 96% detection rate for unseen attacks

4.10.4. Market Basket Analysis (Walmart)

Data: 10M transactions of 5,000 products.

- Mine frequent itemsets with FP-Growth (min support 0.5%)
- Generate association rules with lift > 2
- Discovered that "diapers $\Rightarrow$ beer" in certain evenings (classic example)

4.11. Decision Framework and Recommendations

Choosing the right unsupervised learning method depends on your goal, data characteristics, and constraints. Figure 8 provides a decision flowchart (conceptual; to be inserted as a Mermaid flowchart in final version).

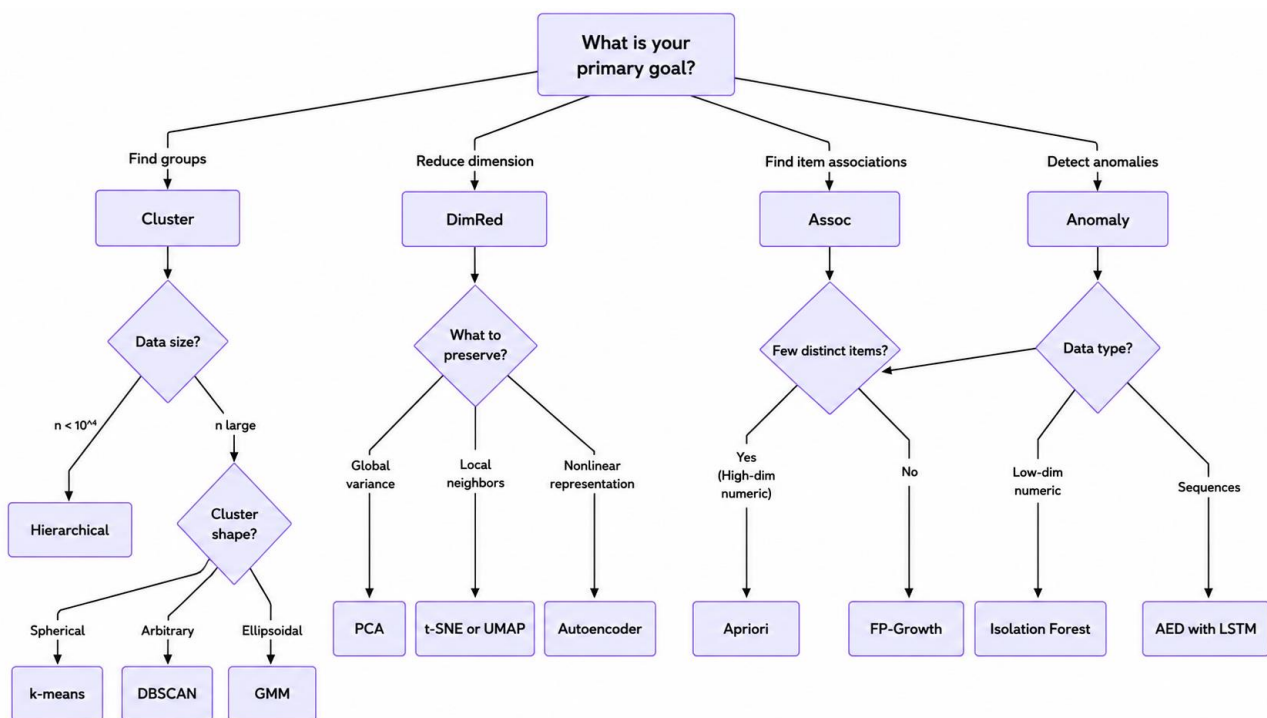


Figure 8: Decision flowchart for selecting unsupervised learning methods based on goal and data characteristics

4.11.1. Recommendations for Practitioners

Table 3 summarizes suitable machine learning methods for different data analysis scenarios and explains the main reasons for selecting each technique.

Table 3: Selection of machine learning methods based on application scenarios

Scenario	Recommended method	Why
Quick exploratory analysis, spherical clusters	k-means with k-means++	Fast, interpretable
Unknown cluster shapes, outliers expected	DBSCAN (tune ϵ)	Robust, finds arbitrary shapes
Hierarchical relationships	Agglomerative clustering with dendrogram	Visual insight
High-dimensional visualization	UMAP (or t-SNE for smaller datasets)	Preserves local structure
Linear feature extraction/noise reduction	PCA	Simple, scalable

Non-linear representation for downstream tasks	Autoencoder	Flexible, can be stacked
Market basket analysis	FP-Growth (or Apriori for small data)	Fast, handles sparse data
Unlabeled anomaly detection	Isolation Forest	Fast, works on high-dim
Complex anomalies (images, sequences)	Autoencoder reconstruction	Captures non-linear patterns
Self-supervised pretraining (images, text)	SimCLR, MAE, BERT	State-of-the-art representations

It gives practical guidance for selecting suitable methods based on data dimensionality, cluster characteristics, computational efficiency, interpretability, and application requirements.

4.11.2. Common Pitfalls

- ✗ Using k-means on non-spherical clusters → misleading results.
- ✗ Interpreting t-SNE distances globally (t-SNE only preserves local structure).
- ✗ Ignoring feature scaling – always standardize/normalize before clustering or PCA.
- ✗ Choosing k by elbow method alone – cross-validate with silhouette.
- ✗ Evaluating clustering without domain validation – external validation metrics (ARI, NMI) are only for benchmarks with labels.

5. Conclusion

Unsupervised learning is an essential toolkit for extracting relevant, useful insights from raw, unlabelled data. Unsupervised methods enable researchers and practitioners to discover hidden structures without extensive manual labeling. Applications include customer segmentation, anomaly detection, visualization of high-dimensional data, and discovery of association rules. Researchers have examined the main approaches using clustering methods such as k-means, DBSCAN, hierarchical clustering, Gaussian mixture models, and spectral clustering; techniques for dimensionality reduction such as PCA, t-SNE, UMAP, and autoencoders; techniques for association mining such as Apriori and FP-Growth; and approaches for anomaly detection based on statistical, proximity, isolation, and reconstruction-based principles. It has also taken into account internal and external evaluation measures, such as silhouette score, Davies-Bouldin index, adjusted Rand index and normalized mutual information, as well as recent advances in self-supervised and contrastive learning. Classical algorithms still provide useful, interpretable, and computationally efficient baselines for exploratory analysis, despite rapid advances. Future work should target scalable learning from streaming data, interpretable deep clustering, robust anomaly detection for high-dimensional spatio-temporal data, and reliable evaluation without human labels. In practice, researchers are encouraged to start with simple baselines, visualize discovered structures, validate their findings with domain expertise, and use modern self-supervised approaches when they have large unlabelled datasets. In the end, unsupervised learning is not a substitute for human judgment, but an augmentation of human ability to discover, investigate, and understand patterns hidden in increasingly complex datasets.

Acknowledgement: The authors gratefully acknowledge Al-Shahbaa Private University, EBLA Private University, Ukrainian National Forestry University, The Fitch Group, Melbourne Institute of Technology, and Messiah University for their valuable support, guidance, research resources, and institutional assistance throughout this research.

Data Availability Statement: The data for this study are available upon request from the corresponding author.

Funding Statement: This manuscript and research paper were prepared without any financial support or funding.

Conflicts of Interest Statement: The authors have no conflicts of interest to declare. This work represents a new contribution by the authors, and all citations and references are appropriately included based on the information utilised.

Ethics and Consent Statement: This research adheres to ethical guidelines and obtains informed consent from all participants. Confidentiality measures were implemented to safeguard participant privacy.

Reference

1. M. Kubsch, C. Krist, and P. Wulff, "Pattern recognition—Unsupervised machine learning," in *Applying Machine Learning in Science Education Research*, Springer, Cham, Switzerland, 2025.

2. R. Agrawal, T. Imieliński, and A. Swami, "Mining association rules between sets of items in large databases," *ACM SIGMOD Record*, vol. 22, no. 2, pp. 207–216, 1993.
3. V. Chandola, A. Banerjee, and V. Kumar, "Anomaly detection: A survey," *ACM Computing Surveys*, vol. 41, no. 3, pp. 1–58, 2009.
4. T. Chen, S. Kornblith, M. Norouzi, and G. Hinton, "A simple framework for contrastive learning of visual representations," in *Proceedings of the 37th International Conference on Machine Learning (ICML)*, Vienna, Austria, 2020.
5. M. Ester, H. P. Kriegel, J. Sander, and X. Xu, "A density-based algorithm for discovering clusters in large spatial databases with noise," in *Proc. of the Second Int. Conf. on Knowledge Discovery and Data Mining (KDD)*, Portland, Oregon, United States of America, 1996.
6. I. J. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial nets," in *Advances in Neural Information Processing Systems*, Montreal, Quebec, Canada, 2014.
7. F. E. Grubbs, "Procedures for detecting outlying observations in samples," *Technometrics*, vol. 11, no. 1, pp. 1–21, 1969.
8. K. He, X. Chen, S. Xie, Y. Li, P. Dollár, and R. Girshick, "Masked autoencoders are scalable vision learners," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, New Orleans, Louisiana, United States of America, 2022.
9. D. P. Kingma and M. Welling, "Auto-encoding variational Bayes," *arXiv preprint arXiv:1312.6114*, 2013. [Accessed by 12/11/2024].
10. S. Lloyd, "Least squares quantization in PCM," *IEEE Transactions on Information Theory*, vol. 28, no. 2, pp. 129–137, 1982.
11. J. MacQueen, "Some methods for classification and analysis of multivariate observations," in *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, Berkeley, California, United States of America, 1967.
12. S. T. Roweis and L. K. Saul, "Nonlinear dimensionality reduction by locally linear embedding," *Science*, vol. 290, no. 5500, pp. 2323–2326, 2000.
13. D. E. Rumelhart, G. E. Hinton, and R. J. Williams, "Learning representations by back-propagating errors," *Nature*, vol. 323, no. 6088, pp. 533–536, 1986.
14. R. R. Sokal and P. H. A. Sneath, "Principles of Numerical Taxonomy," *W. H. Freeman*, San Francisco, California, United States of America, 1963.
15. J. B. Tenenbaum, V. de Silva, and J. C. Langford, "A global geometric framework for nonlinear dimensionality reduction," *Science*, vol. 290, no. 5500, pp. 2319–2323, 2000.
16. L. van der Maaten and G. Hinton, "Visualizing data using t-SNE," *Journal of Machine Learning Research*, vol. 9, no. 86, pp. 2579–2605, 2008.
17. L. McInnes, J. Healy, N. Saul, and L. Großberger, "UMAP: Uniform manifold approximation and projection," *Journal of Open-Source Software*, vol. 3, no. 29, p. 861, 2018.

Publisher's Note: The publisher remains impartial concerning jurisdictional claims in published maps and institutional affiliations. Responsibility for the content rests entirely with the authors and does not necessarily reflect the publisher's perspectives.