

A THEORETICAL REVIEW-LEARNING THEORY

Pawandeep Kaur¹

¹Lecture, Department of BCA, MIMIT, Malout, India.

ABSTRACT

A machine learning program, or simply a learning program, is computer software that learns from experience. A learner program is another name for such a program. There are three sorts of learning theories: supervised, unsupervised, and reinforcement learning. In this chapter, we will look at the three different types of learning theories and compare them based on several criteria.

Keywords: Machine Learning, Supervised Learning, Unsupervised Learning, Reinforcement Learning

1. INTRODUCTION

If a computer program's performance at task T, as measured by P, improves with experience E, the program is said to learn from experience E.

Example: Handwriting recognition learning problem

- Task T: Recognising and classifying handwritten words within images
- Performance P: Percent of words correctly classified
- Training experience E: A dataset of handwritten words with given

The learning process, whether performed by a human or a machine, can be separated into four components: data storage, abstraction, generalization, and evaluation. Figure 1 depicts the many components and processes in the learning process.

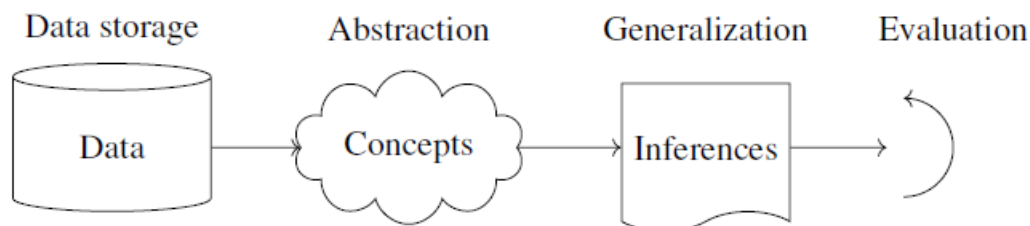


Fig:1. Components and the steps involved in the learning process.

Data storage facilities that can store and retrieve massive volumes of data are a key part of the learning process. Data storage is used as a foundation for advanced reasoning by both humans and machines.

Abstraction is the second component of the learning process. The process of extracting knowledge from recorded data is known as abstraction. This entails developing broad thoughts about the data as a whole. Knowledge production entails both the use of known models and the development of new models.

Generalization is the third component of the learning process. The process of converting knowledge about stored facts into a form that may be used for future action is referred to as generalization. These activities must be carried out on jobs that are similar, but not identical, to those shown previously. In general, the objective is to uncover the data qualities that will be most useful to future activities.

Evaluation: The final step in the learning process is evaluation. It is the process of providing feedback to the user to assess the value of newly acquired information. This input is then used to improve the overall learning process.

2. MACHINE LEARNING

Machine learning is the process of programming computers to maximize a performance criterion based on example data or previous experience. We've defined a model up to some parameters, and learning is the execution of a computer program to optimize the model's parameters using training data or prior experience. The model might be predictive to make future forecasts, descriptive to gather information from data, or both.

Data mining is the application of machine learning technologies to massive databases. A significant number of data is processed in data mining to create a simple model with important applications, such as high predicted accuracy. The following is a list of some of the most common machine learning applications.

- Machine learning is used in retail to study consumer behavior.
- In finance, banks examine historical data to create models for use in loan applications, fraud detection, and stock market trading.

- Learning models are used in production for optimization, control, and troubleshooting.
- Learning programs are utilized in medicine for medical diagnosis.
- In telecommunications, call patterns are examined to optimize networks and maximize service quality.
- Large volumes of data in physics, astronomy, and biology can only be examined quickly enough by computers in science. The World Wide Web is massive; it is continually expanding, and manually looking for relevant information is impossible.
- It is used in artificial intelligence to train a system to learn and adapt to changes so that the system creator does not have to anticipate and supply solutions for every potential circumstance.
- It is used to solve a wide range of issues in vision, speech recognition, and robotics.
- Machine learning approaches are used in the construction of computer-controlled automobiles to ensure proper steering when traveling on a variety of roadways.
- Machine learning approaches have been applied to create programs for games like chess, backgammon, and go.

Types of Learning

In general, machine learning algorithms can be classified into three types.

- Supervised learning
- Unsupervised learning
- Reinforcement learning

2.1 Supervised learning

A training set of examples with the right replies (targets) is supplied, and the algorithm generalizes to reply appropriately to all potential inputs based on this training set. This is also known as learning through examples. The machine learning job of learning a function that translates an input to an output based on example input-output pairs is known as supervised learning.

Each example in the training set in supervised learning is a pair consisting of an input item (usually a vector) and an output value. A supervised learning algorithm examines training data and generates a function that may be used to map fresh samples. The method will properly determine the class labels for unseen instances in the best-case scenario. Both classification and regression problems are examples of supervised learning. There are several supervised learning algorithms available, each with its own set of advantages and disadvantages. There is no single learning algorithm that excels at all supervised learning tasks.

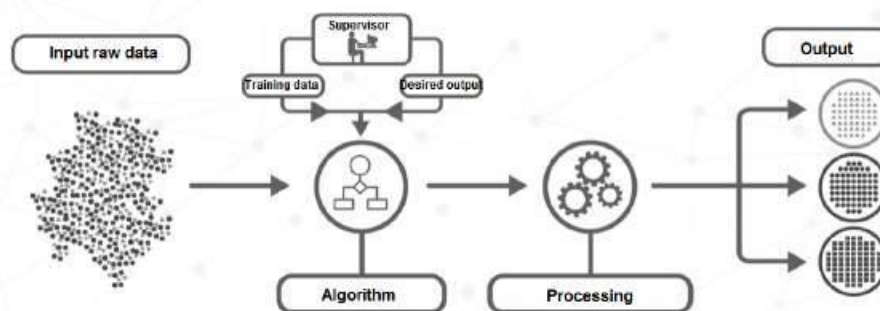


Fig:2. Supervised learning

The term "supervised learning" refers to the process of an algorithm learning from a training dataset that may be compared to a teacher monitoring the learning process. We know the proper answers (or outputs), thus the algorithm generates predictions on the training data repeatedly and is corrected by the teacher. When the algorithm achieves an acceptable level of performance, learning comes to an end.

2.2 Unsupervised learning

Instead of providing correct answers, the algorithm attempts to detect commonalities between the inputs such that inputs with anything in common are grouped. Density estimation is a statistical technique for unsupervised learning.

Unsupervised learning is a sort of machine learning method that is used to derive conclusions from datasets that contain input data but no labeled answers. A classification or categorization is not included in the observations in unsupervised learning methods. There are no output values, hence there is no function estimate. Because the examples provided to the learner are unlabeled, the correctness of the structure produced by the algorithm cannot be assessed. Cluster analysis is the most prevalent unsupervised learning approach, and it is used for exploratory data analysis to uncover hidden patterns or grouping in data.

2.3 Reinforcement learning

This is a hybrid of supervised and unsupervised learning. The algorithm is informed when the answer is incorrect, but not how to correct it. It must investigate and test several alternatives until it determines how to obtain the correct answer. Because the monitor assesses the response but does not advise adjustments, reinforcement learning is also known as learning with a critic.

The task of persuading an agent to act in the environment to maximize its rewards is known as reinforcement learning. A learner (the computer) is not taught what actions to do, as in most kinds of machine learning, but must instead experiment to determine which activities provide the greatest reward. In the most intriguing and demanding scenarios, actions may impact not just the immediate reward but also later situations and, as a result, all subsequent rewards.

3. SUPERVISED LEARNING

A vast variety of supervised learning approaches have been introduced in the field of machine learning over the previous decade. Supervised learning has become a focus of many machine learning research. Many supervised learning approaches have found applications in the processing and analysis of a wide range of data.

One of the key properties of supervised learning is the ability to use annotated training data. In the categorization process, the so-called labels are class. There are several algorithms used in supervised learning approaches. This section covers the key features of a handful of supervised approaches. This review paper's major purpose and contribution are to offer an overview of machine learning as well as machine learning methodologies.

Classification of Supervised Learning Algorithms.

The supervised machine learning methods that deal with classification are as follows: Linear Classifiers, Logistic Regression, Naive Bayes Classifier, Perceptron, Support Vector Machine; Quadratic Classifiers, K-Means Clustering, Boosting, Decision Tree, Random Forest (RF); Neural Networks, Bayesian Networks, and others

1) Linear Classifiers: Linear models for classification use linear (hyperplane) decision boundaries to divide input vectors into classes. The purpose of classification in machine learning linear classifiers is to arrange things with comparable feature values into groups. A linear classifier achieves this purpose, by making a classification decision based on the value of the linear combination of the characteristics.

Because it is the quickest classifier, a linear classifier is frequently employed in cases where classification speed is critical. Also, when the number of dimensions is huge, such as in document classification, where each element is often the number of counts of a word in a document, linear classifiers frequently do extremely well. The rate of convergence among data set variables, on the other hand, is determined by the margin. The margin roughly measures how linearly separable a dataset is, and hence how simple it is to solve a particular classification issue.

2) Logistic regression: This is a classification function that employs a single multinomial logistic regression model with a single estimator and builds on class. In a specific technique, logistic regression generally says where the border between the classes occurs and also states the class probabilities depend on distance from the boundary. When the data set is bigger, this advances faster towards the extremes (0 and 1) These probabilistic claims elevate logistic regression above the level of a simple classifier.

It makes more comprehensive predictions and may be fitted in a different method; yet, such strong forecasts may be incorrect. Logistic regression, like Ordinary Least Squares (OLS) regression, is a prediction method. However, using logistic regression, prediction yields a binary outcome. Logistic regression is a popular method for applied statistics and discrete data processing. Linear interpolation is used in logistic regression.

3) Naive Bayesian (NB) Networks: These are extremely basic Bayesian networks made up of directed acyclic graphs with just one parent (representing the unseen node) and multiple children (corresponding to observed nodes), with a strong assumption of independence among child nodes in the context of their parent. As a result, the independence model (Naive Bayes) relies on estimation.

Bayes classifiers are often less accurate than more advanced learning algorithms (such as ANNs). On standard benchmark datasets, however, performed a large-scale comparison of the naive Bayes classifier with state-of-the-art algorithms for decision tree induction, instance-based learning, and rule induction, and found it to be sometimes superior to the other learning schemes, even on datasets with significant feature dependencies. The attribute independence problem of the Bayes classifier was addressed using Averaged One-Dependence Estimators.

4) Multi-layer Perceptron: This is a classifier in which the network weights are determined by solving a quadratic programming problem with linear constraints rather than a nonconvex, unconstrained minimization problem as in regular neural network training. Other well-known algorithms are based on the perceptron concept. The perceptron algorithm is used to learn from a batch of training examples by repeatedly running the algorithm across the training set

until it finds a prediction vector that is correct across the board. This prediction rule is then applied to the labels on the test set.

5) Support Vector Machines (SVMs): These are the most modern techniques for supervised machine learning. Models of Support Vector Machines (SVMs) are connected to multilayer perceptron neural networks. SVMs are based on the concept of a margin, which is either side of a hyperplane that divides two data classes. It has been demonstrated that increasing the margin and therefore generating the greatest feasible distance between the separating hyperplane and the instances on either side of it reduces the upper bound on the predicted generalization error.

6) K-means: is one of the most basic unsupervised learning techniques for dealing with the well-known clustering problem. The approach follows a basic and uncomplicated method for classifying a given data set using a defined number of clusters (assuming k clusters). When labeled data is unavailable, the K-Means method is used. Converting crude rules of thumb into highly accurate prediction rules is a general technique. With adequate data and a weak learning algorithm that can regularly find classifiers (rules of thumb) that are at least somewhat better than random, say, accuracy Of 55%, a boosting method may provably generate a single classifier with very high accuracy, say, 99%.

7) Decision Trees: Decision Trees (DT) are trees that order instances based on feature values to categorize them. Each node in a decision tree represents a characteristic of a classification instance, and each branch indicates a value that the node might take. Instances are categorized and arranged according to their feature values, beginning with the root node. In data mining and machine learning, decision tree learning employs a decision tree as a prediction model that translates observations about an item into inferences about the item's target value. Classification trees and regression trees are more descriptive names for such tree models. They often use post-pruning approaches to evaluate the performance of decision trees after they have been pruned using a validation set. Any node can be deleted and the most common class of the training instances sorted to it is assigned.

8) Neural Networks: Neural Networks (NN) can execute many regression and/or classification tasks at the same time, even though most networks only perform one. In the great majority of situations, the network will have a single output variable, however, this may equate to several output units in the case of many-state classification issues (the post-processing stage takes care of the mapping from output units to output variables). The Artificial Neural Network (ANN) is based on three key aspects: the unit's input and activation functions, network design, and the weight of each input link. Given that the first two characteristics are fixed, the ANN's behavior is governed by the current weight values.

9) Bayesian Network: A Bayesian Network (BN) is a graphical model that depicts the probability correlations between a collection of variables (features). The most well-known statistical learning methods are Bayesian networks. When compared to decision trees or neural networks, the most intriguing characteristic of BNs is undoubtedly the ability to take into account previous information about a particular problem in terms of structural correlations among its elements. One limitation of BN classifiers is that they are not ideal for datasets with a large number of characteristics.

Supervised machine learning techniques are useful in a wide range of fields. In general, SVMs and neural networks outperform when dealing with multidimensional and continuous information. When dealing with discrete/categorical characteristics, logic-based systems tend to perform better. A high sample size is necessary for neural network models and SVMs to attain maximum prediction accuracy, whereas NB may use a very small dataset. There is widespread agreement that k -NN is extremely sensitive to irrelevant characteristics; this trait may be explained by the algorithm's operation. Furthermore, the existence of irrelevant information might make neural network training ineffective, if not impossible. Most decision tree algorithms struggle to solve issues that require diagonal partitioning. The instance space is divided orthogonally to one variable's axis and parallels the other axes. As a result, the partitioned areas are all hyperrectangles. When there is multi-collinearity and a nonlinear relationship between the input and output features, ANNs and SVMs work well. During both the training and classification stages, Naive Bayes (NB) requires little storage space: the strict minimum is the memory required to hold the prior and conditional probabilities. The basic k NN method requires a large amount of storage space during the training phase, and its execution space is at least as large as its training space. On the contrary, because the final classifier is generally a highly condensed summary of the data, execution space is frequently significantly lower than training space for all non-lazy learners. Furthermore, although rule algorithms cannot be utilized as incremental learners, Naive Bayes and k NN can. Missing values are naturally ignored in computing probability and so have no influence on the final choice in Naive Bayes. k NN and neural networks, on the other hand, require entire records to function. Finally, Decision Trees and NB have distinct operating characteristics; while one is very accurate, the other is not, and vice versa. Decision trees and rule classifiers, on the other hand, have a similar operational profile. SVM and ANN both have comparable operating profiles. Overall datasets, no single learning algorithm can consistently outperform other algorithms.

4. UNSUPERVISED LEARNING

Unsupervised learning, as the name implies, is a machine learning approach in which models are not supervised by training datasets. Instead, models discover hidden patterns and insights in the provided data. It is comparable to the learning that occurs in the human brain while learning new things. It is defined as "a kind of machine learning in which models are trained on unlabeled datasets and are then allowed to operate on that data without supervision."

Because, unlike supervised learning, we have input data but no corresponding output data, unsupervised learning cannot be applied directly to a regression or classification task. The purpose of unsupervised learning is to discover the underlying structure of a dataset, categorize it based on similarities, and display it compactly. Working of unsupervised learning can be understood by the below diagram:

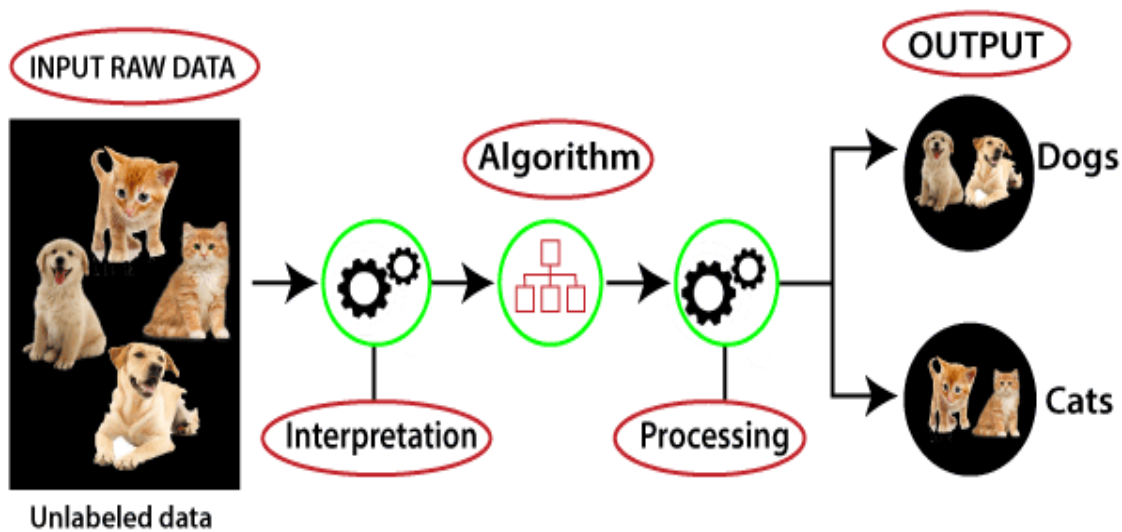


Fig:3. Working of unsupervised learning algorithm

We used unlabeled input data in this case, which means it is not classified and no matching outputs are provided. This unlabeled input data is now supplied into the machine learning model to train it. It will first analyze the raw data to uncover hidden patterns in the data and then use appropriate algorithms such as k-means clustering, Decision tree, and so on. Once the appropriate method is applied, the algorithm splits the data items into groups based on their similarities and differences.

Types of Unsupervised Learning Algorithm:

The unsupervised learning algorithm can be further categorized into two types of problems:

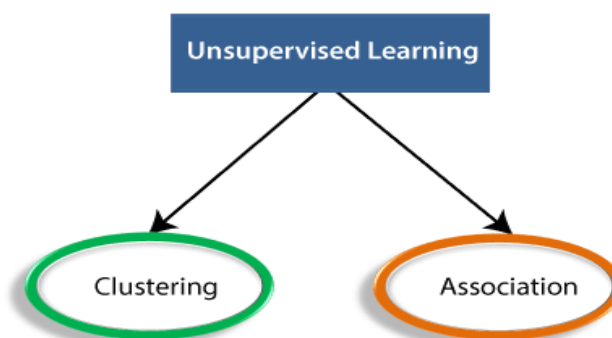


Fig:4. Types of unsupervised learning

Clustering is a way of organizing things into clusters so that objects with the highest similarities stay in one group while having little or no similarities with objects in another group. Cluster analysis discovers similarities between data items and categorizes them based on the existence or absence of such similarities.

An association rule is a type of unsupervised learning strategy used to discover links between variables in a big database. It determines the collection of elements in the dataset that appear together. The association rule improves the effectiveness of the marketing strategy. People who purchase X (say, bread) are more likely to purchase Y (Butter/Jam). Market Basket Analysis is a common example of an Association rule.

Unsupervised learning algorithms:

K-means is a clustering method that is also known as partitioning or segmentation. It divides the data points into a predetermined number of clusters known as K. In the K-means method, K is the input since you tell the computer how many clusters you want to identify in your data. Each data item is subsequently assigned to the closest cluster center, known as a centroid (black dots in the picture). The latter serve as data storage spaces.

The Fuzzy K-means algorithm is a modification of the K-means algorithm, which is used to do overlapping clustering. Unlike the K-means method, fuzzy K-means assumes that data points might belong to many clusters with varying degrees of proximity to each. The distance from a data point to the cluster's centroid is used to calculate proximity. As a result, there may be some overlap across various clusters.

Gaussian Mixture Models (GMMs) are a probabilistic clustering technique. Because the mean and variance are unknown, the models assume a fixed number of Gaussian distributions, each representing a distinct cluster. The procedure is used to determine which cluster a certain data point belongs to.

Hierarchical clustering technique: Each data point may be given to a different cluster in the hierarchical clustering technique. The two clusters that are closest to each other are then combined to form a single cluster. The merging process is repeated until only one cluster remains at the top. Bottom-up or agglomerative approaches are examples of such approaches. Top-down or divisive hierarchical clustering is the process of starting with all data items tied to the same cluster and then performing splits until each data item is assigned as a distinct cluster.

Recommender systems. The association rules approach is used for analyzing buyer baskets and detecting cross-category purchase relationships. Amazon's "Frequently bought together" suggestions are a wonderful example. The firm intends to develop more successful up-selling and cross-selling methods, as well as make product recommendations based on the frequency of specific goods discovered in one shopping cart.

Target marketing. Regardless of industry, the approach of association rules may be utilized to extract rules to aid in the development of more successful target marketing strategies. A travel agency, for example, may utilize customer demographic information as well as historical data from prior campaigns to determine which categories of consumers to target for their new marketing campaign.

The Apriori algorithm uses common itemsets to construct association rules. The things with the highest support value are called frequent itemsets. The program builds itemsets and discovers relationships by repeatedly scanning the whole dataset.

The principal component analysis is an algorithm used for dimensionality reduction. It is used to minimize the number of features within huge datasets, resulting in greater data simplicity without sacrificing accuracy. The feature extraction procedure is used to compress datasets. It indicates that elements from the original set are joined to form a new, smaller one. These new characteristics are referred to as principal components.

5. REINFORCEMENT LEARNING

Reinforcement learning (RL) is a broad technique for problem-solving that is focused on rewards. RL attempts to emulate how humans learn new things by interacting with their surroundings rather than through a teacher. For example, when a newborn learns to wave hands, cry, and laugh, it is influenced by parental input. When we learn to drive a car, we learn to turn left and right to prevent a collision on the road. RL is the process through which robots learn to attain a goal through interactions with their surroundings. RL is also a sequential decision-making and control issue in mathematics. For example, when driving a car, we must decide whether to turn left or right after making the preceding decision.

Reinforcement learning entails performing appropriate actions to maximize reward in a specific scenario. It is used by various applications and computers to determine the best potential action or course to take in a given scenario. Reinforcement learning differs from supervised learning in that the training data contains the solution key, allowing the model to be trained with the right answer, but in reinforcement learning, there is no answer and the reinforcement agent determines what to do to complete the given task. It is obligated to learn from its experience in the absence of a training dataset. One sort of machine learning is reinforcement learning. Supervised learning is the most well-known sort of machine learning. Algorithms are built in supervised learning to produce outputs that mirror the labels supplied in the training set. In contrast to supervised learning, providing a supervisor in the problem of RL is challenging since we frequently have no notion of what the proper option is. For example, if we wish to drive a car, we cannot identify every photo taken by the camera. RL is widely employed in engineering, neurology, psychology, mathematics, and economics, in addition to driving an automobile.

The dilemma is as follows: we have an agent and a reward, but there are several obstacles in between. The agent is expected to locate the shortest way to the prize. The difficulty is better explained in the following problem.

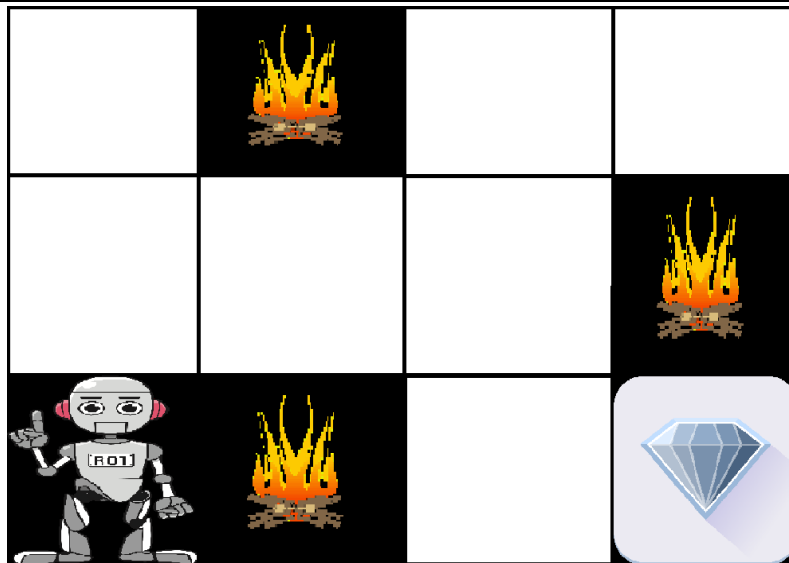


Fig:5. Example to understand reinforcement learning

The graphic above depicts the robot, diamond, and fire. The robot's purpose is to obtain the diamond prize while avoiding the obstacles that are launched. The robot learns by attempting all potential pathways and then selecting the option that provides the best reward with the fewest obstacles. Each correct step awards the robot a reward, while each incorrect step deducts the robot's payout. When it reaches the ultimate prize, the diamond, the entire reward will be computed.

Main points in Reinforcement learning

Input: The input should be a starting state from which the model will begin.

Output: There are several possible outputs since there are numerous solutions to a certain problem.

Training: The model will return a state, and the user will select whether to reward or penalize the model depending on its output.

The model is always learning.

The optimal answer is determined by the highest possible reward.

Types of Reinforcement:

There are two types of Reinforcement:

Positive

Positive reinforcement happens when an event that occurs as a result of a certain behavior improves the strength and frequency of the behavior. In other words, it influences conduct positively. The following are the benefits of reinforcement learning:

Maximizes performance and allows for long-term change.

Excessively Reinforcement can result in an excess of states, which can reduce the effectiveness of the outcomes.

Negative

Bad Reinforcement is defined as behavior strengthening as a result of a negative circumstance being ended or avoided. Benefits of reinforcement learning:

Improves Behavior Provide defiance to a minimum performance standard It just supplies enough to satisfy the bare minimum of conduct.

Various Practical applications of Reinforcement Learning are:

In robotics, RL may be used for industrial automation. Machine learning and data processing may both benefit from RL. RL may be used to construct training systems that deliver students with customized education and resources based on their needs.

In vast areas, RL may be employed in the following scenarios:

The environment has a model, but there is no analytical solution;

There is only a simulation model of the environment provided (the subject of simulation-based optimization)

Interacting with the environment is the only method to gather knowledge about it.

Comparing Supervised, unsupervised, and reinforcement learning

Key distinctions between supervised, unsupervised, and reinforcement learning:

Regression and classification are the two fundamental goals of Supervised Learning. Clustering and associative rule mining difficulties are addressed via unsupervised learning. Whereas Reinforcement Learning is concerned with exploitation or exploration, Markov's decision processes, Policy Learning, Deep Learning, and value learning are concerned with learning.

Supervised Learning works with labeled data, and the machine is aware of the output data patterns. However, unsupervised learning works with unlabeled data and produces results based on a collection of impressions. In contrast, in the Markov Decision Process of Reinforcement Learning, the agent interacts with the environment in discrete stages. The name itself says, Supervised Learning is highly supervised.

And Unsupervised Learning is not supervised., Reinforcement Learning is less supervised which depends on the agent in determining the output.

In Supervised Learning, the input data is labeled data. Unsupervised Learning, on the other hand, uses unlabeled data. Reinforcement Learning does not need preexisting data. Supervised Learning makes predictions depending on the kind of class. Unsupervised Learning uncovers hidden patterns. The learning agent also serves as a reward and action mechanism in Reinforcement Learning.

Comparison Table:1

Criteria	Supervised Learning	Unsupervised Learning	Reinforcement Learning
Definition	Learns by using labeled data	Trained using unlabelled data without any guidance	Works on interacting with the environment
Type of data	Labeled data	Unlabelled data	No-predefined data
Type of problems	Regression and classification	Association and clustering	Exploitation or exploration
Supervision	Extra supervision	No supervision	No supervision
Aim	Calculate outcomes	Discover underlying patterns	Learn a series of action
Application	Risk evaluation, Forecast sales	Recommendation system, Anomaly Detection	Self-Driving cars, Gaming, Healthcare

6. CONCLUSION

The amount of data produced in the globe nowadays is enormous. This information is created not just by humans, but also by cell phones, computers, and other electronic devices.

A programmer would undoubtedly determine how to train an algorithm utilizing a certain learning model based on the type of data accessible and the motive offered.

Machine Learning is a branch of computer science in which a system's performance improves through executing tasks repeatedly using data rather than explicitly designed by programmers. Let us now distinguish between three Machine Learning techniques: supervised, unsupervised, and reinforcement learning.

Consider yourself a student in a school where your teacher is watching you to see "how you can solve the problem" or "if you are doing it right or not." Similarly, when Supervised Learning input is presented as a labeled dataset, a model may readily learn from it to deliver the issue outcome.

Unsupervised learning, on the other hand, lacks a comprehensive and clean labeled dataset. Self-organized learning is unsupervised learning. Its primary goal is to investigate the underlying patterns and anticipate the outcome. We simply provide the computer data and tell it to hunt for hidden characteristics and logically cluster the data.

However, reinforcement learning is not dependent on either supervised or unsupervised learning. Furthermore, in this case, the algorithms learn to react to their surroundings on their own. It is quickly expanding and creating a wide range of learning algorithms.

These algorithms are useful in robotics, gaming, and other fields.

Labeled data is mapped to known output via supervised learning. Unsupervised Learning, on the other hand, investigates patterns and predicts outcomes. Reinforcement Learning is a trial-and-error process. To summarise, the purpose of Supervised Learning is to build formulas based on input and output data. We identify a relationship between input values and group them in Unsupervised Learning. In Reinforcement Learning, an agent learns by interacting with the environment and receiving delayed feedback.

7. REFERENCES

- [1] ZHOU, Z. H. I.-H. U. A. (2022). Machine learning. SPRINGER VERLAG, SINGAPOR.
- [2] Machine Learning, statistics, and Data Analytics. (2021). Machine Learning. <https://doi.org/10.7551/mitpress/13811.003.0005>
- [3] Characterizing simulated traffic management initiatives with unsupervised learning. (2022). <https://doi.org/10.2514/6.2022-4025.vid>
- [4] Detailed study of unsupervised machine learning clustering efficacy in identifying unstable appro... (2022). <https://doi.org/10.2514/6.2022-3968.vid>
- [5] Young, R. and Ringenberg, J. (2019). Machine Learning: An Introductory Unit of Study for Secondary Education. Proc. of the 50th ACM Technical Symposium on Computer Science Education, New York, NY, USA, 1274
- [6] Voulgari, I., et al. (2021). Learn to Machine Learn: Designing a Game Based Approach for Teaching Machine Learning to Primary and Secondary Education Students. Interaction Design and Children. New York, NY, USA, 593–598.
- [7] Wan, X., et al. (2020). SmileyCluster: supporting accessible machine learning in K-12 scientific discovery. Proc. of the Interaction Design and Children Conference. New York, NY, USA
- [8] Zhu, K. (2019). An Educational Approach to Machine Learning with Mobile Applications. M.Eng thesis Elect. Eng. Comput. Sci., Massachusetts Institute of Technology, Cambridge, MA, USA. Zimmermann-Niefield, A., et al. (2019).
- [9] Youth Learning Machine Learning through Building Models of Athletic Moves. Proc. of the 18th ACM Int.Conference on Interaction Design and Children, Boise, ID, USA.
- [10] Tedre, M., et al. (2021). Teaching Machine Learning in K–12 Classroom: Pedagogical and Technological Trajectories for Artificial Intelligence Education. IEEE Access, (9), 110558-110572.
- [11] Mike, K. and Rosenberg-Kima. R. (2021). Teaching Machine Learning to Computer Science Preservice Teachers: Human vs. Machine Learning. Proc. of the 52nd ACM Technical Symposium on Computer Science Education, New York, NY, USA, 1368
- [12] McBride, E., et al.. (2021). Design Considerations for Inclusive AI Curriculum Materials. Proc. of the 52nd ACM Technical Symposium on Computer Science Education, New York, NY, USA, 1370.
- [13] Micheuz, P. (2020). Approaches to Artificial Intelligence as a Subject in School Education. IFIP Advances in Information and Communication Technology, Open Conference on Computers in Education, Mumbai, India
- [14] Kaspersen, M. H., et al. (2021) The Machine Learning Machine: A Tangible User Interface for Teaching Machine Learning. Proc. of the 15th International Conference on Tangible, Embedded, and Embodied Interaction, New York, NY, USA, Article 19, 1–12.
- [15] Joshua, J. (2021). Integrating Machine Learning in Secondary Geometry, Mathematics Teacher: Learning and Teaching PK-12, 114(4), 325-329.
- [16] Huang, Chao-Jung., et al. (2021). Developing a medical artificial intelligence course for high school students. Proc. of the International Forum on Medical Imaging in Asia, Taipei, Taiwan, 11792.
- [17] Gresse von Wangenheim, C., et al. (2021). Visual tools for teaching machine learning in K-12: A ten-year systematic mapping. Education and Information Technology, 26, 5733–5778.
- [18] L. Samuel, “Some studies in machine learning using the game of checkers,” IBM Journal of Research and Development, vol. 3, pp. 210–229, 1959.