



A Comprehensive Review of Machine Learning Frameworks and Practical Applications

Ravi Bhushan, Vineeta Khemchandani



Abstract: A fundamental component of digitalisation solutions that have garnered significant attention in the digital sphere is machine learning, which is primarily an area of artificial intelligence. The author's goal in this study is to provide a concise overview of the most widely utilised machine learning algorithms for this purpose. To assist in selecting the best learning algorithm to meet the application's specific needs, the author aims to highlight the advantages and limitations of machine learning algorithms from the standpoint of their application. This paper provides a brief overview and outlook on the various uses of machine learning techniques.

Keywords: Machine Learning, Pseudo Code, Algorithm, Supervised Learning, Unsupervised Learning, Reinforcement Learning, ML Application and Task

Nomenclature:

SVM: Support Vector Machine

ML: Machine Learning

TSVMs: Transductive Support Vector Machines

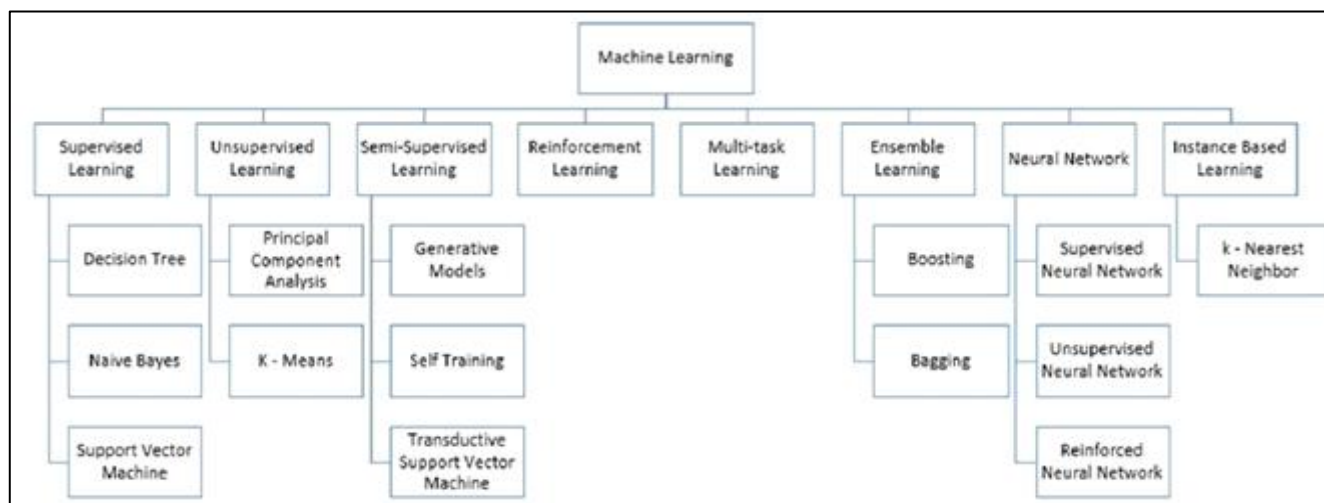
MTL: Multi-Task Learning

KNN: K-Nearest Neighbours

I. INTRODUCTION

Humans have used a wide range of tools to improve their

performance in various jobs since the beginning of time. The human brain's inventiveness created many machines. These gadgets make life easier for humans by meeting a variety of demands, such as those related to industry, transportation, and computing. Machine learning [1] is also one of them. Arthur Samuel defined machine learning as the study of teaching computers to learn without explicit programming. Arthur Samuel became well-known for his checkers program. The process of teaching machines [2] to handle data more efficiently is known as machine learning (ML). After analysing the data, we might discover that we are unable to interpret it. We employ machine learning when that happens. Machine learning is becoming increasingly necessary as more and more datasets become available. Many industries utilise machine learning to extract relevant data. Machine learning aims to learn from data. The topic of teaching robots to learn without explicit programming has been extensively studied. Numerous mathematicians and programmers employ a range of methods to address this issue, including those that involve enormous datasets.



[Fig.1: Type of Machine Learning]

Manuscript received on 02 March 2026 | Revised Manuscript received on 09 March 2026 | Manuscript Accepted on 15 March 2026 | Manuscript published on 30 March 2026.

*Correspondence Author(s)

Ravi Bhushan*, School of Computer Science and Engineering. Galgotias University, Gr. Noida (Uttar Pradesh), India. Email ID: Ravidelhi2524@gmail.com, ORCID ID: 0009-0001-7224-9169

Dr. Vineeta Khemchandani, School of Computer Applications and Technology, Galgotias University, Gr. Noida (Uttar Pradesh), India. Email ID: Vineetakh05@gmail.com, ORCID ID: 0000-0002-7934-8493

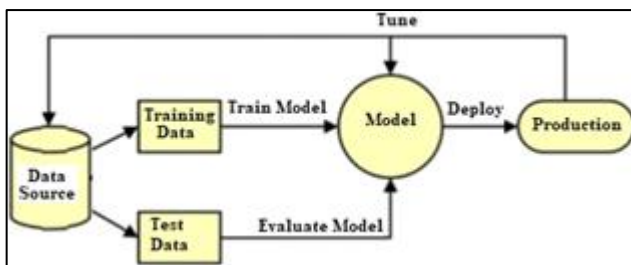
© The Authors. Published by Blue Eyes Intelligence Engineering and Sciences Publication (BEIESP). This is an open-access article under the CC-BY-NC-ND license <http://creativecommons.org/licenses/by-nc-nd/4.0/>

Machine learning tackles data problems using a range of methods. Data scientists would like to stress that there isn't a single, all-encompassing approach that is effective in all circumstances. The kind of problem you wish to solve, the number of variables, the model that would be most appropriate for it, and other considerations all influence the kind of approach that is employed. Here is [3] a quick summary of some of the most widely used machine learning (ML) algorithms.



II. SUPERVISED LEARNING

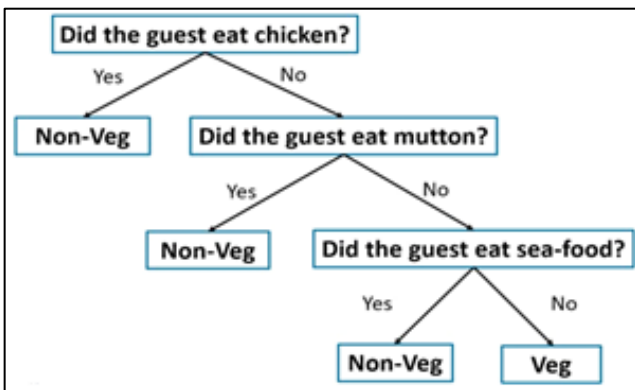
Supervised learning is a machine learning task in which sample input-output pairs are used to learn a function that maps inputs to outputs. To infer a function, it uses labelled training data and a set of training samples. Supervised algorithms are machine learning algorithms that need external assistance. The input dataset contains both the train and test datasets. The output variable in the training dataset must be projected or categorised. All algorithms look for patterns in the training dataset and use them to classify or predict on the test dataset. The figure below shows the workflow of a supervised machine learning algorithm. The most popular supervised machine learning techniques have been discussed in this article [7].



[Fig.2: Supervised learning Workflow]

A. Decision Tree

A decision tree is a graph that shows options and their outcomes. The nodes in the graph represent events or options, while the edges represent circumstances or decision rules. Every tree contains nodes and branches. Each branch denotes a value that the node can accept, and each node indicates the qualities of a group that require classification.



[Fig.3: Decision Tree]

B. Decision Tree Pseudo Code:

```

def
Decision Tree Learning (examples, attributes,
parent_examples):
if len(examples) == 0:
return plurality Value (parent_ examples)
# Return the most probable answer, as there is no training
data left
elif len(attributes) == 0:
return plurality Value (examples)
elif (all examples classify the same):
return their classification
A = max (attributes, key(a)=importance (a,
examples) # choose the most promising attribute to
    
```

```

condition on tree = new Tree(root=A)
for value in A.values():
exs = examples [e.A == value]
subtree = decision Tree Learning (exs, attributes. Remove
(A), examples)
# Note implementation should probably wrap the trivial case
returns into trees for consistency tree. add Subtree as Branch
(subtree, label= (A, value)
return tree
    
```

C. Naive Bayes

This classification approach, which assumes predictor independence, is based on Bayes' Theorem. In short, a Naive Bayes classifier asserts that the presence of a certain trait inside a class is unrelated to the presence of any other attribute. The text classification industry is the main market for Naïve Bayes. Based on the conditional likelihood of occurrence, its main uses are in classification and grouping.

$$P(c|x) = \frac{P(x|c)P(c)}{P(x)}$$

Likelihood
Class Prior Probability

Posterior Probability
Predictor Prior Probability

$$P(c|X) = P(x_1|c) \times P(x_2|c) \times \dots \times P(x_n|c) \times P(c)$$

[Fig.4: Naive Bayes]

Pseudo Code of Naive

Bayes Input:

Training dataset T,

F = (f1, f2, f3..., fn) // value of the predictor variable in the testing dataset.

Output: A class of testing datasets.

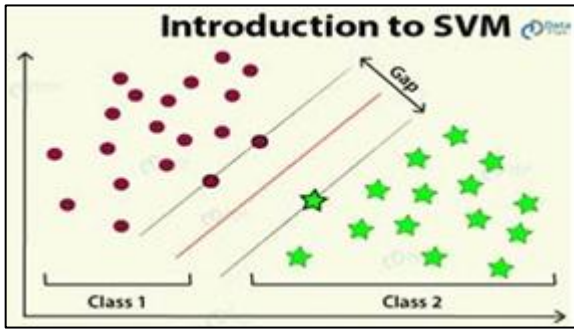
Steps:

- 1) Read the training dataset T;
- 2) Calculate the mean and standard deviation of the predictor variables in each class;
- 3) Repeat, calculate the probability of fi using the Gauss density equation in each class, until the probability of all predictor variables (f1, f2, f3,..., fn) has been calculated.
- 4) Calculate the likelihood for each class.
- 5) Get the greatest likelihood

D. Support Vector Machine

Support Vector Machine (SVM) is the most widely used state-of-the-art machine learning technique. In machine learning, support vector machines are supervised learning models with associated learning algorithms that assess data for classification and regression. SVMs can perform non-linear classification in addition to linear classification thanks to the kernel trick, which implicitly maps inputs to high-dimensional feature spaces. The main concept is creating boundaries between the classes. The margins are drawn to maximise the distance between the margin and the classes, thereby decreasing classification error.





[Fig.5: Support Vector Machine]

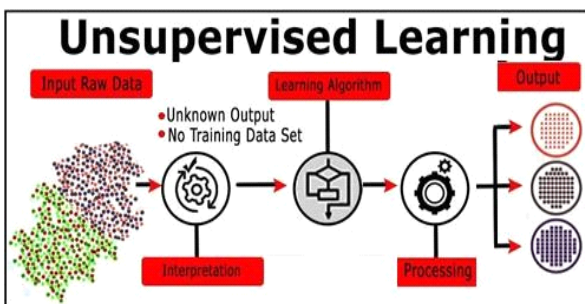
Pseudo Code of Support Vector Machine

```

initialise  $Y_i = YI$  for  $i \in I$ 
repeat
    compute SVM solution  $vv, b$  for the data set with imputed labels
    compute outputs  $ii = (vv, xi) + b$  for all  $xi$  in positive bags
    set  $yi = \text{sgn}(fi)$  for every  $i \in i$ ,  $yi = 1$ 
    for (every positive bag  $bi$ )
    end
    if  $(liei(1 + yi)/2 == 0)$ 
    compute  $i^* = \arg \max_{i \in i} ii$ 
    set  $yi^* = 1$ 
    end
    while (imputed labels have changed)
    output ( $vv, b$ )
    
```

III. UNSUPERVISED LEARNING

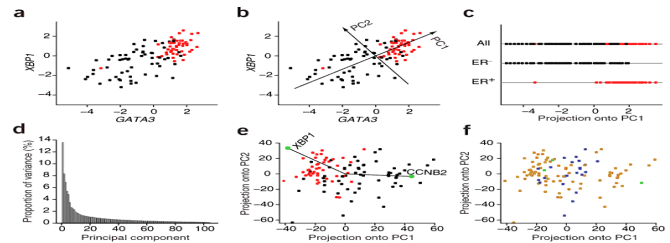
These are called unsupervised learning because, unlike supervised learning described above, there are no correct answers or tutors. Algorithms are left to discover and present the fascinating structure in the data. Only a few features are identified by the unsupervised learning algorithms from the data. It uses the previously learnt features to identify the class of newly presented data. Its main uses are in feature reduction and clustering.



[Fig.6: Unsupervised Learning]

A. Principal Component Analysis

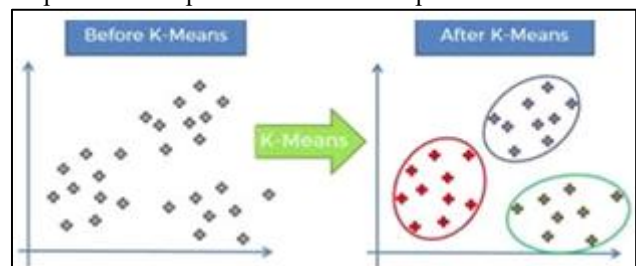
Principal component analysis is a statistical technique that creates principal components—a collection of values of linearly uncorrelated variables—from a set of observations of potentially correlated variables using an orthogonal transformation. To speed up and simplify the computations, the data's dimensionality is reduced. It uses linear combinations to explain the variance-covariance structure of a set of variables. It is frequently applied as a method of dimensionality reduction.



[Fig.7: Principal Component Analysis]

B. K-Means Clustering

The well-known clustering problem can be solved by one of the most straightforward unsupervised learning techniques, K-means. The process uses a straightforward method to group a given data set into a predetermined number of clusters. Define k centres, one for each cluster, as the primary idea. It is important to situate these centres, as different locations yield varied strategic outcomes. Therefore, it is preferable to position them as far apart as feasible.



[Fig.8: K-Means Clustering]

Pseudo Code of K-Means Clustering

```

Input:
 $D = \{t_1, t_2, \dots, t_n\}$  // Set of elements
 $K$  // Number of desired clusters
Output:
 $K$  // Set of clusters
K-Means algorithm:
    Assign initial values for  $m_1, m_2, \dots, m_k$ 
    repeat
        assign each item  $t_j$  to the clusters which has the closest mean; calculate new mean for each cluster;
    until convergence criteria is met;
    
```

IV. SEMI-SUPERVISED LEARNING

Unsupervised and supervised machine learning techniques are combined to create semi-supervised machine learning. In fields like data mining and machine learning, when unlabelled data is already available and obtaining labelled data is laborious, it may be beneficial. A machine learning algorithm is trained on a "labelled" dataset, where each record contains the result information, in increasingly popular supervised machine learning techniques. The following section discusses a few semi-supervised learning algorithms.

A. Transductive SVM

Semisupervised learning has extensively used transductive support vector machines (TSVMs) to handle partially labelled data. Because its basis in generalisation is not well understood, there has been a mystery about it. It is employed to label the unlabelled data so that the

margin between the labelled and unlabelled data is as small as possible. Using TSVM to find an exact answer is NP-hard.

B. Generative Models

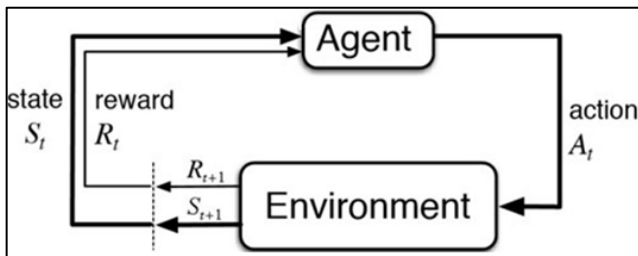
A model that can produce data is called a generative model. Both the class (i.e., the entire data) and the features are modelled. Since I can create data points using this probability distribution if we model $P(x,y)$, all algorithms that model $P(x,y)$ are generative. For each component, one labelled example is sufficient to verify the mixture's distribution.

C. Self-Training

Self-training involves using a subset of labelled data to train a classifier. Unlabelled data is then supplied to the classifier. In the training set, the unlabelled points and the anticipated labels are totalled. After that, the same process is repeated. The term "self-training" refers to the classifier learning from itself.

V. REINFORCEMENT LEARNING

The study of how software agents should behave in a given environment to maximise a concept of cumulative reward is known as reinforcement learning. The three fundamental machine learning paradigms are reinforcement learning, supervised learning, and unsupervised learning.



[Fig.9: Reinforcement Learning]

VI. MULTITASK LEARNING

A type of machine learning called "multi-task learning" seeks to solve several tasks simultaneously by leveraging their similarities. This can serve as a regularizer, increasing the effectiveness of learning. In a formal sense, Multi-Task Learning (MTL) will improve the learning of a particular model by utilising the knowledge contained in all n tasks. This is because traditional deep learning approaches aim to solve only one task using a single model. This is the case if there are n tasks, or a subset of them, that are related to each other but not identical.

VII. ENSEMBLE LEARNING

Ensemble learning is the process of intentionally generating and combining multiple models, such as classifiers or experts, to address a specific computational intelligence problem. Ensemble learning is most often used to improve a model's performance or to reduce the likelihood that an inferior model is selected by accident. Other uses of ensemble learning include feature selection, data fusion, incremental learning, non-stationary learning, error-correcting, and assigning a confidence level to the model's choice.

A. Boosting

"Boosting" describes a class of algorithms that transform weak learners into strong ones. Boosting is an ensemble learning strategy that reduces variance and bias. The foundation of boosting is the question posed by Kearns and Valiant: "Can a set of weak learners create a single strong learner?" A classifier that has an arbitrarily high correlation with the true classification is considered a strong learner. In contrast, a weak learner is a classifier with low correlation with the true classification.

B. Pseudo Code of Boosting

Boosting Algorithm under Loss Function L
Inputs: Joint distribution Q on $\mathcal{Z} = \mathcal{X} \times \{1, -1\}$ and initial predictor $H^{(0)} \in \mathcal{S}[\mathcal{H}]$. In practice, Q is the empirical probability of training samples.
Loop For $m = 1, \dots, M$
 1. Find a hypothesis $h^{(m)} \in \mathcal{H}$ such that

$$h^{(m)} = \arg \min_{h \in \mathcal{H}} \int_{\mathcal{Z}} Q(dz) L'(-y H^{(m-1)}(x)) I(y \neq h(x))$$
 The minimization does not need to be exact.
 2. Find a coefficient $\alpha^{(m)} \in \mathbb{R}$ such that

$$\alpha^{(m)} = \arg \min R_L(Q, H^{(m-1)} + \alpha h^{(m)}).$$
 Line search methods or Newton methods can be applied to solve the one-dimensional optimization problem.
 3. Update the predictor: $H^{(m)} = H^{(m-1)} + \alpha^{(m)} h^{(m)}.$
Output: $H^{(M)}$ as an estimated predictor.

C. Bagging

In situations where improving the accuracy and stability of a machine learning algorithm is necessary, bagging (bootstrap aggregating) is used. Both regression and classification can use it. Additionally, bagging helps manage overfitting and reduces variance.

D. Pseudo Code of Bagging

```

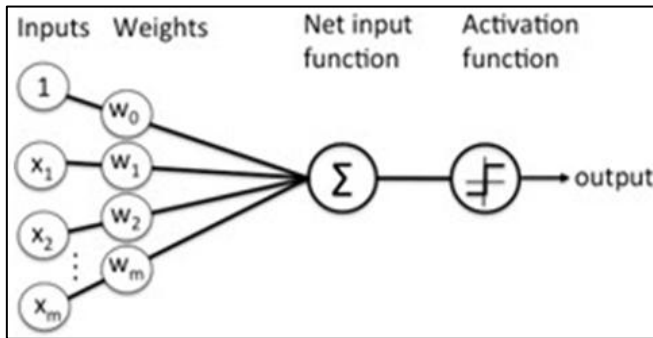
1. Initialize feature vector bg feature= [0, 0, ..., 0]
2. for token in text.tokenize() do
3.     if token in dict then
4.         token idx=getindex(dict, token)
5.         bg feature[token idx]++
6.     else
7.         continue
8.     end if
9. end for
10. return bg feature
    
```

VIII. NEURAL NETWORKS

To identify underlying links in a set of data, a neural network is a collection of algorithms that simulates how the human brain functions. Neural networks, in this context, are networks of neurons, whether they be artificial or natural. Because they can adapt to changing inputs, neural networks produce optimal results without changing the output criteria. As trading systems are developed, the idea of neural networks—which has its



origins in artificial intelligence—is rapidly gaining traction.

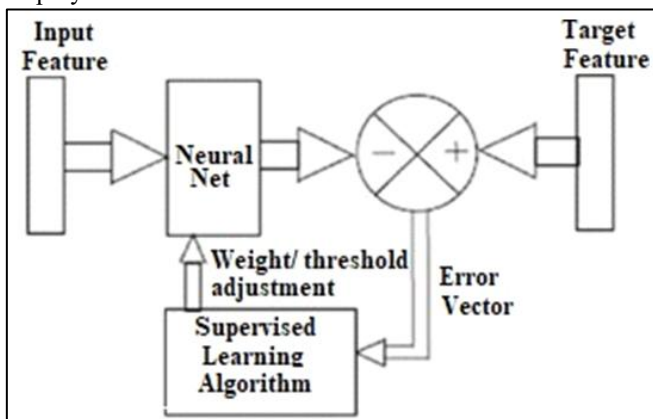


[Fig.10: Neural Networks]

In the same way, an artificial neural network acts. It operates on three levels. The input layer receives input. Within the concealed layer, the input is processed. The calculated output is finally sent via the output layer.

A. Supervised Neural Network

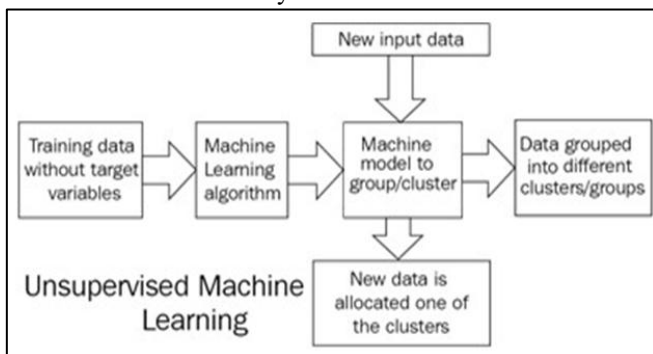
In a supervised neural network, the output is predetermined by the input. The neural network's actual output and its anticipated output are contrasted. The parameters are adjusted and then re-fed into the neural network based on the error. In feedforward neural networks, supervised learning is employed.



[Fig.11: Supervised Neural Network]

B. Unsupervised Neural Network

The input and output are unknown to the neural network beforehand. Classifying data by similarity is the network's primary function. After grouping different inputs, the neural network determines if they are correlated.

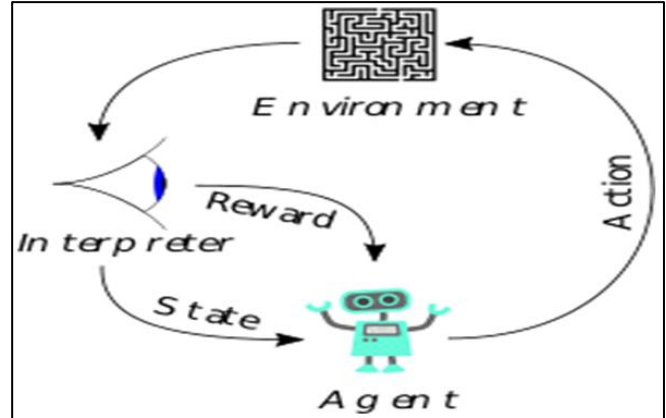


[Fig.12: Unsupervised Neural Network]

C. Reinforced Neural Network

Reinforcement learning is the term used to describe goal-oriented algorithms that learn to maximise a specific metric

over a large number of steps, for example, maximising the number of points earned in a game over a large number of moves, or achieving a difficult aim (goal). They can begin with nothing and, given the correct circumstances, perform at superhuman levels. These algorithms are reinforced when they make the correct choices and punished when they make the wrong ones, much like a toddler who is rewarded with candy and spankings.



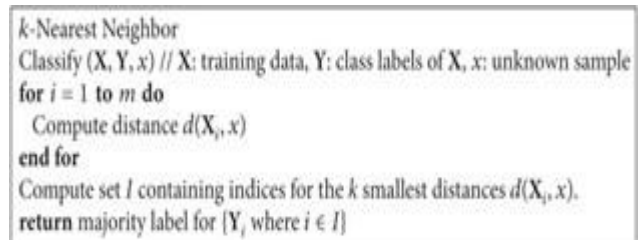
[Fig.13: Reinforced Neural Network]

IX. INSTANCE-BASED LEARNING

A class label or prediction is generated via instance-based learning, a family of classification and regression techniques, based on how similar the query is to its nearest neighbour (s) in the training set. Instance-based learning algorithms explicitly differ from other approaches, like decision trees and neural networks, in that they do not generate an abstraction from individual instances. Instead, they merely store all the information and obtain an answer at query time by analysing the query's nearest-neighbour(s).

A. K-Nearest Neighbour

One easy supervised machine learning method for handling regression and classification issues is the k-nearest neighbours (KNN) algorithm. It is simple to use and comprehend, but it has a key flaw: it becomes slower as the amount of data used increases.



X. MACHINE LEARNING APPLICATION AND TASK.

Statistical arbitrage, learning associations, classification, prediction, medical diagnosis, digital image processing (image recognition) [5], large data analysis [4], speech recognition, and many other applications employ machine learning techniques [6].

Generally, machine learning tasks fall into one of three major groups based on the type of "signal" or "feedback" that a learning system can use. It is the machine learning task of

inferring a function from labelled training data. One has to carry out the following steps: (i) Decide on the kind of training examples. (ii) Collect a training set. (iii) Decide the input (iv) Decide the structure of the learned function and the corresponding learning algorithm. (v) Complete the design. (vi) Run the learning algorithm on the gathered training set. (vii) Assess the accuracy of the learned function.

XI. CONCLUSION

One can choose between supervised and unsupervised machine learning. Supervised learning is the better option if you have fewer data and clearly marked training data. For large datasets, unsupervised learning typically performs better and yields better outcomes. Choose deep learning methods if you have a large data collection at your disposal. Additionally, you have studied Deep Reinforcement Learning and Reinforcement Learning. Neural networks, their uses, and their limits are now clear to you. This study examines several machine learning algorithms. Nowadays, whether they realise it or not, everyone uses machine learning. using social networking sites to update images or receive product recommendations when buying online. This publication provides an overview of the most well-known machine learning algorithms.

DECLARATION STATEMENT

Some of the cited references are older and are noted explicitly as [7]. However, these works remain significant for the current study, as they are pioneering in their fields.

As the article's author, I must verify the accuracy of the following information after aggregating input from all authors.

- **Conflicts of Interest/ Competing Interests:** Based on my understanding, this article has no conflicts of interest.
- **Funding Support:** This article has not been funded by any organizations or agencies. This independence ensures that the research is conducted objectively and without external influence.
- **Ethical Approval and Consent to Participate:** The content of this article does not necessitate ethical approval or consent to participate with supporting documentation.
- **Data Access Statement and Material Availability:** The adequate resources of this article are publicly accessible.
- **Author's Contributions:** The authorship of this article is contributed equally to all participating individuals.

REFERENCES

1. Batta Mahesh, "Machine Learning Algorithms -A Review", International Journal of Science and Research (IJSR), 2019. DOI: <http://doi.org/10.21275/ART20203995>
2. Susmita Ray, "A Quick Review of Machine Learning Algorithms", International Conference on Machine Learning, Big Data, Cloud and Parallel Computing (Com-IT-Con), India, 14th -16th Feb 2019. DOI: <https://doi.org/10.1109/COMITCon.2019.8862451>
3. Rushika Ghadge, Juilee Kulkarni, Pooja More, Sachee Nene, Priya R, "Prediction of Crop Yield using Machine Learning", International Research Journal of Engineering & Technology, Vol 5, Issue 2, Feb-2018. <https://www.irjet.net/archives/V5/i2/IRJET-V5I2479.pdf>
4. Sonal S. Ambalkar, S. S. Thorat2, "Bone Tumour Detection from MRI Images using Machine Learning: A Review", International Research Journal of Engineering & Technology, Vol. 5, Issue 1, Jan -2018. <https://www.irjet.net/archives/V5/i1/IRJET-V5I1112.pdf>
5. Diksha Sharma and Neeraj Kumar, "A Review on Machine Learning

Algorithms, Tasks and Applications", International Journal of Advanced Research in Computer Engineering & Technology (IJARCET), Volume 6, Issue 10, October 2017, ISSN: 2278 – 1323. <https://www.researchgate.net/publication/320609700>

6. Kumar, N. and Gupta, S., 2016. Offline Handwritten Gurmukhi Character Recognition: A Review. International Journal of Software Engineering and Its Applications, 10(5), pp.77-86. DOI: <https://doi.org/10.14257/ijseia.2016.10.5.08>
7. Muhammad, I. and Yan, Z., 2015. Supervised Machine Learning Approaches: A Survey. ICTACT Journal on Soft Computing, 5(3). <https://www.ictactjournals.in/IJSC/ArticleDetails?id=1785>, works remain significant, see the [declaration](#)

AUTHOR'S PROFILE



Mr. Ravi Bhuhan is a PhD scholar (School of Computer Science and Engineering) at Galgotias University, Uttar Pradesh, India. He has over 17 research papers in various international journals and peer-reviewed conferences. He has authored a chapter on Artificial Intelligence and Machine Learning in two books. He is the Editor-in-Chief of Arts and Communication Journals.



Dr. Vineeta Khemchandani is the Dean (School of Computer Applications and Technology) at Galgotias University, Uttar Pradesh, India. She has over 50 research papers in various international journals and peer-reviewed conferences. She has received the 50 Best Education Leaders award from the World Education Congress and the Most Active Participation (Woman) award from the Computer Society of India. She has served as the Guest Editor of several special issues in various Scopus- and Emerging Sources Citation Index (Clarivate Analytics)-indexed conferences and journals. She has authored and reviewed several books on Operating Systems, Data structures and Algorithms, and UNIX programming. She is the Editor of PLOS Journals, Journal of King Saud University, Science Publishing Group USA, American Association of Science and Technology, and Journal of Electronic Government, an International Journal, INDERSCIENCE Publishers. Dr Vineeta is one of the developers of various research-based solutions in Cognitive Science, including Cognitive profiling for recruitment, Drone surveillance, marine simulators, and mental health detectors. She has received research grants from various national funding agencies in India.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of the Blue Eyes Intelligence Engineering and Sciences Publication (BEIESP)/ journal and/or the editor(s). The Blue Eyes Intelligence Engineering and Sciences Publication (BEIESP) and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.