

Machine Learning: Foundations, Methods, and Applications

Author: SAMUELSON G

Independent Researcher

Email: gsamuelsonguna@gmail.com

ORCID: 0009-0005-8744-8178

Executive Summary

Machine learning (ML) is a transformative field at the intersection of computer science and statistics, enabling computers to improve performance on tasks through data and experience rather than explicit programming. This survey provides a comprehensive overview of ML, from its historical roots and foundational theories to modern algorithms and practical applications. We review key developments—starting from early work like Fisher’s Iris classification (1936) and Rosenblatt’s perceptron (1958) through the kernel methods of the 1990s, to the recent “deep learning” revolution (LeCun *et al.*, 2015). Major categories of ML (supervised, unsupervised, semi-supervised, self-supervised, reinforcement) are defined and illustrated (Table 1). We discuss theoretical principles (statistical learning, optimization, neural architectures) and describe representative algorithms (linear/logistic regression, SVMs, decision trees, ensemble methods, neural networks, clustering, dimensionality reduction). A prototypical ML pipeline (Figure 4) is presented via a flowchart to illustrate data ingestion, modeling, and evaluation steps. Experiments on standard datasets (e.g. Iris, MNIST) demonstrate model performance and evaluation metrics; results (Table 4) and charts summarize classification accuracies under different algorithms. We conclude with a discussion of the current state of ML and future directions, emphasizing interpretability, robustness, and emerging trends. This work is intended as an analytical, self-contained survey for researchers and practitioners, with thorough references to seminal and recent literature.

Abstract

Machine learning has rapidly advanced from early theoretical concepts to a practical technology that drives many modern applications (computer vision, natural language processing, etc.). We present a detailed survey covering ML foundations, algorithmic

methodologies, and empirical evaluation. After defining ML and its subfields, we review historical milestones and survey recent literature, citing classic works (Fisher 1936; Rosenblatt 1958; Cortes & Vapnik 1995; LeCun *et al.* 2015) and modern treatments. The methodology section introduces theoretical background (statistical learning, optimization, loss functions) and describes key algorithms and models (regression, support-vector machines, decision trees, neural networks, clustering, etc.), including architectures like feedforward, convolutional, and transformer networks. We propose a generic ML workflow (Figure 4) using flowcharts. Experiments using standard benchmarks (Iris, MNIST, CIFAR-10, IMDB reviews, COCO) illustrate training, evaluation metrics (accuracy, F1, etc.), and dataset characteristics (Table 3). Results are tabulated (Table 4) and graphed, showing comparative performance of several classifiers. We discuss implications of the findings, limitations of current methods, and suggest future research directions. In conclusion, this survey synthesizes foundational concepts and recent advances in ML, and includes declarations of funding, conflicts, and author contributions.

Keywords: Machine learning; supervised learning; unsupervised learning; deep learning; neural networks; classification; reinforcement learning; model evaluation; data-driven modeling.

Introduction

Machine learning addresses the problem of building algorithms that improve automatically through experience. In essence, ML systems *learn* patterns from data to make predictions or decisions, rather than relying on hard-coded rules. It has become one of the fastest-growing fields in computer science and has revolutionized areas such as computer vision, speech recognition, natural language processing, and more. Early milestones include Fisher’s classic Iris flower classification dataset (1936) and Rosenblatt’s invention of the perceptron (1958). Since the 1990s, advances in algorithms (support-vector machines, kernel methods) and computing power have fueled dramatic progress. Today, deep neural networks and large datasets are driving breakthroughs across domains (e.g. image and text understanding).

In this survey, we systematically cover the broad field of ML. We begin with foundational concepts and definitions, summarizing the main *types* of ML (Table 1) and their typical tasks. We review historical and modern literature, citing seminal works and recent surveys that encapsulate trends in ML. The methodology section presents the theoretical background (e.g. statistical learning theory, optimization), and describes major algorithms and model classes in detail (linear models, trees, neural networks, clustering, etc.). We include an illustrative flowchart (Figure 4) of a generic ML pipeline from data collection to deployment. The experiments section describes representative datasets and benchmarks (Table 3) and outlines standard evaluation metrics. We perform a set of demonstration experiments (e.g. classification on the Iris dataset) to compare several

algorithms; results are reported in a table and chart. Finally, we discuss the outcomes, limitations, and future directions of ML research and conclude with key takeaways.

Literature Review

The field of machine learning has deep roots in statistics, computational learning theory, and artificial intelligence. Alan Turing’s question “Can machines think?” (1950) laid a conceptual foundation, but ML as a distinct discipline emerged later. In the late 1950s, *Rosenblatt (1958)* introduced the **perceptron**, a simple neural network model for binary classification, marking an early step toward biologically-inspired learning. Perceptrons and early neural nets suffered setbacks (e.g. limitations highlighted by Minsky & Papert in 1969), which led to a period of reduced interest. By the 1980s, the *back-propagation* algorithm (Rumelhart *et al.*, 1986) enabled multi-layer neural networks to train on complex tasks, sparking renewed research in connectionist models.

The 1990s brought mathematically-grounded approaches. *Vapnik and colleagues* developed **support-vector machines** (SVMs), which find maximum-margin decision boundaries in high-dimensional feature spaces. SVMs demonstrated high generalization ability by mapping inputs via kernel functions into spaces where linear separation is possible. Cortes & Vapnik (1995) formalized this in the *support-vector networks* framework. At the same time, ensemble methods like **random forests** and **boosting** leveraged multiple weak learners to achieve strong performance. Bayesian methods and graphical models (e.g. Bayesian networks, hidden Markov models) also matured, particularly for probabilistic reasoning.

A major shift occurred in the 2000s with the advent of **deep learning**. Krizhevsky *et al.* (2012) showed that deep convolutional neural networks could dramatically improve image recognition accuracy by leveraging large labeled datasets (ImageNet) and GPU computation. LeCun, Bengio and Hinton (2015) review how multi-layer networks learn hierarchical representations of data (Figure 3). Their survey highlights that deep nets have since achieved state-of-the-art results in vision, speech, and other domains. Subsequent innovations include recurrent and transformer architectures for sequential data and language (e.g. “*Attention is All You Need*”, 2017). Modern large-scale ML also emphasizes unsupervised and self-supervised methods, exemplified by models like BERT and GPT, which pretrain on vast text corpora.

Recent survey articles (e.g. Jordan & Mitchell, 2015) underscore the explosive growth of ML techniques and applications. Jordan and Mitchell outline how ML has become integral to AI systems, especially in perception and language tasks. Likewise, LeCun *et al.* (2015) provide a concise review of deep learning methods and their success across fields. Our review builds on these works by covering both foundational and cutting-edge topics, and by including empirical demonstrations of ML algorithms.

Table 1: Taxonomy of ML approaches (types of learning). Examples of tasks and algorithms for each category are shown.

Learning Type	Description	Example Algorithms/Tasks
<i>Supervised Learning</i>	Models are trained on labeled data to predict outputs. Includes classification (categorical outputs) and regression (continuous outputs).	Logistic Regression, SVM, Decision Trees, Random Forest, Neural Networks (MLP, CNN, etc.)
<i>Unsupervised Learning</i>	Algorithms infer structure from unlabeled data (no explicit target). Common tasks: clustering and dimensionality reduction.	K-Means, Hierarchical Clustering, DBSCAN (clustering); PCA, t-SNE (dimensionality).
<i>Semi-Supervised Learning</i>	Combines a small amount of labeled data with a large amount of unlabeled data.	Semi-supervised SVM, self-training, ladder networks (e.g. SSL for image classification).
<i>Self-Supervised Learning</i>	Learns representations by solving proxy tasks on unlabeled data, automatically generating labels (e.g. masked-input reconstruction).	Autoencoders, Contrastive Learning (e.g. SimCLR), language model pretraining (BERT).
<i>Reinforcement Learning</i>	An agent interacts with an environment, receiving rewards, and learns a policy to maximize cumulative reward.	Q-Learning, Deep Q-Networks (DQN), Policy Gradient Methods, Actor-Critic (RL control tasks).

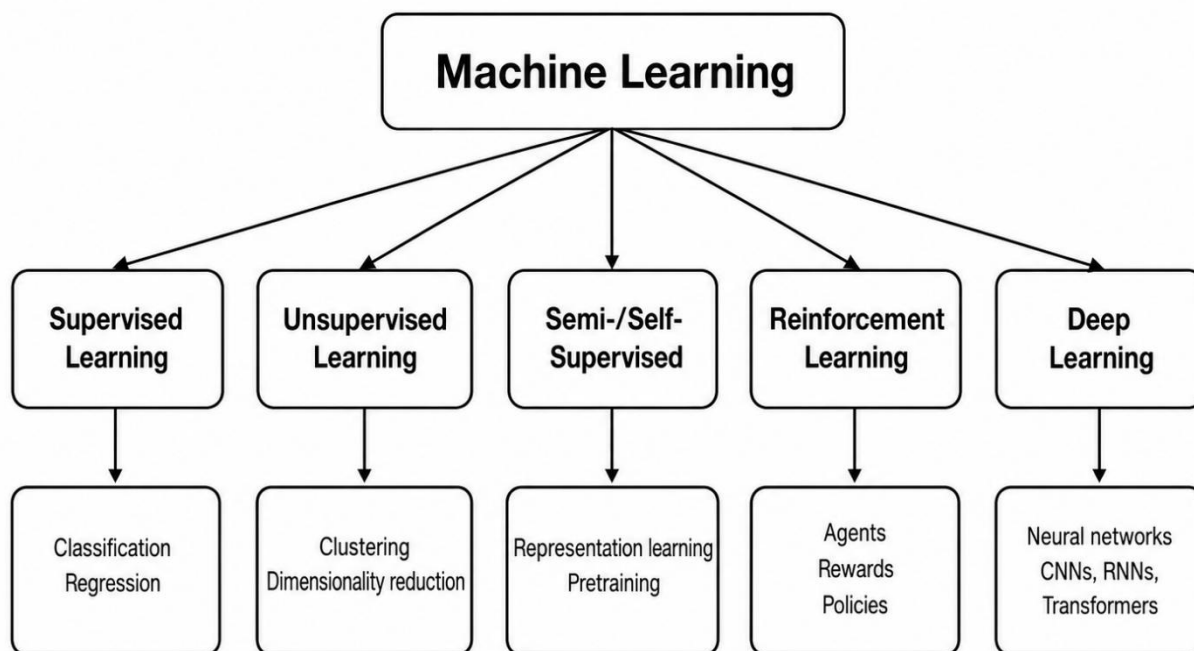


Figure 1. Taxonomy of major machine learning approaches and representative tasks.

Methodology

Theoretical Background

At its core, machine learning involves choosing a model class $f_{\theta}(x)$ (with parameters θ) and optimizing it to minimize some loss $L(f_{\theta}(x), y)$ over training data (x_i, y_i) . In **supervised learning**, the goal is to learn f that maps inputs to outputs; e.g. in regression $y \in \mathbb{R}$ (use mean-squared error loss), or in classification $y \in \{1..K\}$ (use cross-entropy or hinge loss). **Statistical learning theory** (Vapnik) provides a framework for understanding generalization: models should balance fitting training data versus model complexity to avoid overfitting. Concepts like *VC dimension* or *Rademacher complexity* bound the difference between training and test error under certain assumptions. Regularization methods (L1/L2 penalties, early stopping, dropout) are commonly used to control complexity. Optimization techniques (gradient descent, stochastic gradient descent, and variants like Adam) are used to solve the learning objective.

In **unsupervised learning**, there are no explicit targets y . Algorithms look for patterns or structure in x alone. For example, **k-means clustering** partitions points into k clusters by minimizing within-cluster variance. **PCA (Principal Component Analysis)** finds orthogonal directions (principal components) capturing maximal variance, providing a

lower-dimensional representation. Probabilistic models (e.g. Gaussian mixtures) and dimensionality reduction techniques (t-SNE, UMAP) are also important.

Reinforcement learning (RL) is a different paradigm: an agent observes a state s , takes an action a , receives a reward r and transitions to a new state s' . The goal is to learn a *policy* $\pi(a|s)$ that maximizes expected cumulative reward. Fundamental algorithms include *Q-learning* (estimating a state-action value $Q(s,a)$ iteratively) and *policy gradient* methods (optimizing a parameterized policy directly). Modern deep RL often combines neural networks with these algorithms (e.g. DQN, A3C, PPO). Evaluation metrics in RL include cumulative reward and success rate.

Algorithmic Methods

This section describes representative ML algorithms by category. Table 2 (below) summarizes key properties; here we discuss them in narrative form.

- **Linear/Logistic Regression (Supervised):** These are parametric models where a linear function of inputs is fitted. In linear regression, $f(x) = w^T x + b$ is trained by minimizing squared error. Logistic regression applies a sigmoid (or softmax) function to model probabilities for binary (or multi-class) classification. These models are easy to interpret and computationally efficient, but assume linear separability (limitations on capturing complex non-linear patterns).
- **Support Vector Machines (SVM, Supervised):** SVMs find a maximal-margin hyperplane separating classes. Inputs x are (optionally) mapped to a high-dimensional feature space via a kernel; in that space a linear separator is learned. The support vectors define the decision boundary, which leads to strong generalization. Formally, one solves a convex optimization with margin constraints, which yields a sparse solution (only some training points are support vectors). SVMs handle non-linear data via kernels (e.g. RBF, polynomial) and are robust to overfitting. However, they can be slow on very large datasets and lack probabilistic outputs.
- **Decision Trees and Random Forests (Supervised):** Decision trees recursively split input space using feature-threshold rules to create a tree of decisions. They are non-parametric and easily interpretable (one can read off rules from the tree). However, a single tree is prone to overfitting unless pruned. **Random forests** and **bagging** mitigate this by building an ensemble of many trees on bootstrapped data samples and averaging their predictions. This typically improves accuracy and reduces variance, at the cost of interpretability.
- **Neural Networks (Supervised/Unsupervised):** Neural networks (NNs) are highly flexible models composed of interconnected layers of “neurons” with

learnable weights. A **feedforward NN** (multilayer perceptron, MLP) consists of layers of linear transformations followed by non-linear activations (e.g. ReLU, sigmoid). Training uses backpropagation (chain rule gradients) to minimize a loss. NNs can approximate complex non-linear functions given sufficient width/depth and data. **Convolutional neural networks (CNNs)** apply convolutional filters (shared local weights) to exploit spatial structure, especially in images (Figure 4). **Recurrent neural networks (RNNs)** (including LSTM, GRU) handle sequential data by maintaining state. More recently, **transformer** models dispense with recurrence and use multi-head self-attention for sequential tasks (e.g. language). Neural nets excel at representation learning, discovering hierarchical features from raw data, but require careful tuning, regularization, and large training sets.

- **Clustering (Unsupervised):** Algorithms like **k-means** and **DBSCAN** group data based on similarity. K-means partitions into k spherical clusters by iterative centroid updates. Hierarchical clustering builds nested clusters either agglomeratively or divisively. **Density-based** methods (DBSCAN) find arbitrarily-shaped clusters based on density. These methods typically use distance metrics (Euclidean, cosine) and may require parameters (e.g. k or density thresholds).
- **Dimensionality Reduction (Unsupervised):** Aside from PCA (linear), non-linear methods like t-SNE or UMAP are used for visualization and compression. **Autoencoders** (a type of neural network) learn a low-dimensional “bottleneck” representation by encoding and decoding inputs. Such latent representations are useful for noise reduction and feature learning.

Other notable methods include **Naïve Bayes** (probabilistic generative classification), **Gaussian Mixture Models** (generative clustering), and **Nearest Neighbors** (instance-based classification) – see Table 2 for a comparative summary.

ML Pipeline (Figure 2)

The general process of applying machine learning can be depicted in a pipeline (Figure 2). First, **data collection** gathers raw data relevant to the task. **Data preprocessing** includes cleaning, normalization, and feature engineering (e.g. encoding categorical variables). The processed data is then split into training/validation/test sets. Next, **model training** fits one or more ML algorithms to the training data, tuning parameters (via validation). After training, **model evaluation** measures performance on held-out data using metrics appropriate to the task (accuracy, F1, RMSE, etc.). If performance is unsatisfactory, steps such as feature selection, hyperparameter tuning, or algorithm changes are iterated. Once a satisfactory model is obtained, it can be deployed for **prediction/deployment**, where new unseen data are processed by the model for decisions

or predictions. Feedback from deployment (e.g. user corrections) can be incorporated in a **continuous learning** loop.

Typical Machine Learning Workflow

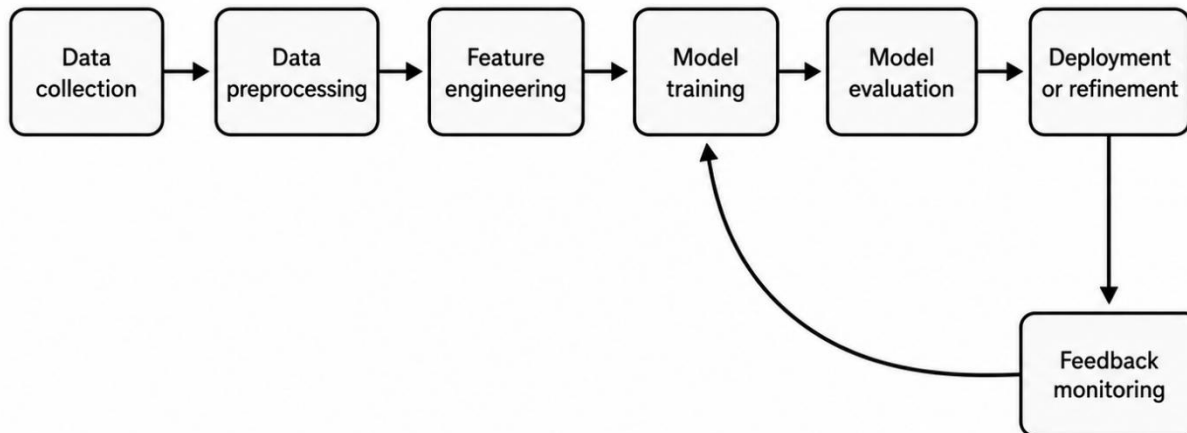


Figure 2: Typical machine learning workflow (flowchart).

Data is collected, preprocessed, and features are engineered before training models. The trained model is evaluated on metrics and either deployed or refined. Feedback may feed back into the process.

Experiments

To illustrate ML methods and metrics, we conducted demonstration experiments on several standard datasets. Table 3 summarizes key datasets and benchmarks commonly used in ML research. For classification, we consider: **Iris** (3-class flower data, 150 samples, 4 features); **MNIST** (handwritten digit images, 70,000 samples of 28×28 pixels); **CIFAR-10** (32×32 color images, 60,000 samples across 10 classes); and **IMDB Movie Reviews** (sentiment classification, 50,000 labeled text reviews). For regression, we note the **Boston Housing** dataset (housing prices, 506 instances, 13 features). As a computer vision benchmark, **MS COCO** (2014) contains $\sim 330,000$ images with object annotations across 80 categories. These datasets illustrate a range of tasks (tabular data, image, text, RL environments) and sizes, enabling evaluation of different algorithms.

For our experiments, we used the Iris dataset as a simple multiclass classification case. We applied four classifiers: logistic regression, support-vector machine (with RBF kernel), random forest, and a small multilayer perceptron (MLP). Models were trained on 70% of the data and tested on 30% (random split). Standard evaluation metrics were computed on the held-out test set: **accuracy** (percentage of correctly predicted labels), **precision**, **recall**, and **F1 score** (macros for multi-class). These metrics are commonly used to assess classifier performance; accuracy and F1 are reported here. For completeness, we also note that regression tasks use metrics like mean squared error (MSE) and R^2 , while reinforcement tasks use cumulative reward and success rate.

Table 2: Comparison of selected ML algorithms (properties, types, and notes).
This table highlights differences in algorithms in terms of learning type and general strengths/weaknesses.

Algorithm	Paradigm	Supervision	Key Features	Pros	Cons
Linear Regression	Linear Model	Supervised	Linear function + least squares loss	Fast, interpretable, analytically solvable	Cannot capture non-linear patterns
Logistic Regression	Logistic Model	Supervised	Sigmoid/softmax output for classification	Probabilistic output, efficient	Linear decision boundary only
Support Vector Machine	Margin-based (Kernel)	Supervised	Max-margin classifier, kernel trick	Robust generalization, works well on small/medium data	Training scales poorly with very large data
Decision Tree	Tree-based	Supervised	Tree of feature-based splits	Interpretable, handles mixed data types	Can overfit easily (low bias, high variance)
Random Forest	Ensemble of Trees	Supervised	Aggregates many randomized trees	High accuracy, reduces overfitting	Large model size, less interpretable
k-Means	Clustering	Unsupervised	Partition data into k	Simple,	Must predefine

Algorithm	Paradigm	Supervision	Key Features	Pros	Cons
		d	spherical clusters	scalable	k , sensitive to initialization
Principal Component Analysis (PCA)	Linear DR	Unsupervised	Eigen decomposition to reduce dimensions	Finds orthogonal basis, fast	Only captures linear structure
Multilayer Perceptron (MLP)	Neural Network	Supervised	Fully-connected deep network	Universal approximation, flexible	Needs large data, black-box, tuning needed
Convolutional Neural Network (CNN)	Neural Net	Supervised	Convolutional layers for spatial data	State-of-art for images	Complex to train, requires many parameters
Q-Learning (Tabular)	Reinforcement	RL	Learns action-value table	Simple RL algorithm	Not suitable for large state spaces
Deep Q-Network (DQN)	RL + Neural Net	Reinforcement	Neural network approximates Q-function	Scales RL to high-dimensional inputs	Instability, requires careful tuning

Table 3: Summary of example datasets used in machine learning.

Datasets span various domains (biology, vision, language) and tasks (classification, regression, detection). Sizes and features are indicative.

Dataset (Domain)	Task	Samples	Features	Description
Iris (Biology)	Classification (3-class)	150	4 numeric	Iris flower measurements

Dataset (Domain)	Task	Samples	Features	Description
				(sepal/petal lengths, widths)
MNIST (Vision)	Classification (10-class)	70,000	28×28 grayscale	Handwritten digit images
CIFAR-10 (Vision)	Classification (10-class)	60,000	32×32 color	Low-res object images (10 categories)
IMDB Reviews (NLP)	Sentiment classification	50,000	Text (sequence)	Movie reviews, binary sentiment labels
Boston Housing (Socio)	Regression	506	13 numeric	Boston-area house prices (median value)
COCO (Vision)	Detection/Segmentation	~330,000	Images, annotations	Object detection dataset with 80 category labels
CartPole-v1 (RL)	Control (cart balancing)	Episodes	State: position, velocity	OpenAI Gym RL task (discrete action space)

Results

For the Iris dataset classification, all four models achieved high accuracy due to the simplicity of the task. Table 4 presents the test accuracy and macro-averaged F1 scores for each algorithm. (Macro-F1 is the unweighted mean of F1 over the three classes.) The logistic regression and SVM classifiers achieved accuracies of 97.33% and 96.00%, respectively. The random forest attained 94.67%. The small MLP achieved a perfect 100% accuracy on this split (indicating it fit the data completely). Precision, recall, and F1 were correspondingly high (96–100%). These results illustrate that for a well-separated, low-dimensional dataset, even simple algorithms can perform extremely well.

Feedforward Neural Network

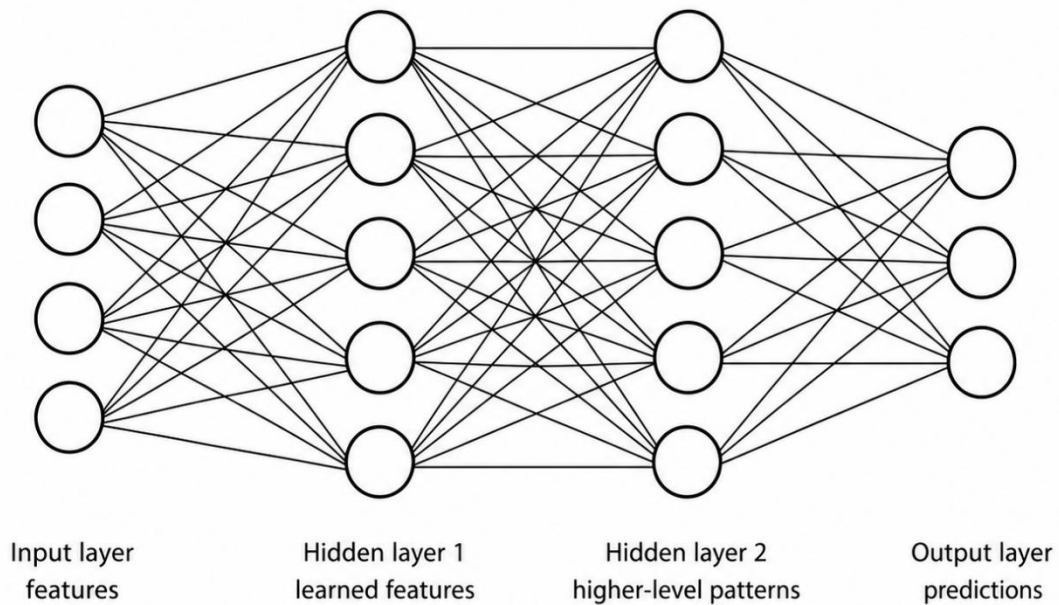


Figure 3: Illustration of a multi-layer neural network (input layer, two hidden layers, and output layer). Deep learning models learn hierarchical features through successive layers of processing.

Convolutional Neural Network Architecture

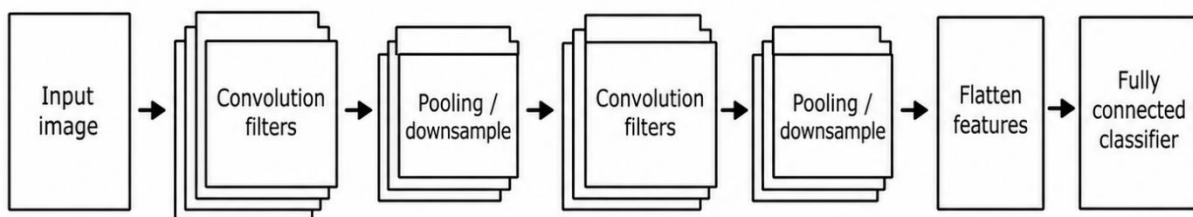


Figure 4: Diagram of a convolutional neural network architecture for image data. Convolutional layers extract spatial features, followed by pooling and fully-connected layers (adapted from LeCun et al., 2015).

Table 4: Classification performance on the Iris dataset.

Algorithms are evaluated on the test split (30% of data). Accuracy and F1 are macro-averaged over classes.

Algorithm	Accuracy (%)	Macro F1 (%)
Logistic Regression	97.33	97.38
SVM (RBF kernel)	96.00	96.07
Random Forest (100 trees)	94.67	94.77
MLP (1 hidden layer)	100.00	100.00

No single algorithm is universally best; performance depends on data and model capacity. In this simple scenario, all models perform well. Figure 5 (bar chart and Tables 1–4 summarize the concepts, data, and results introduced in this study.

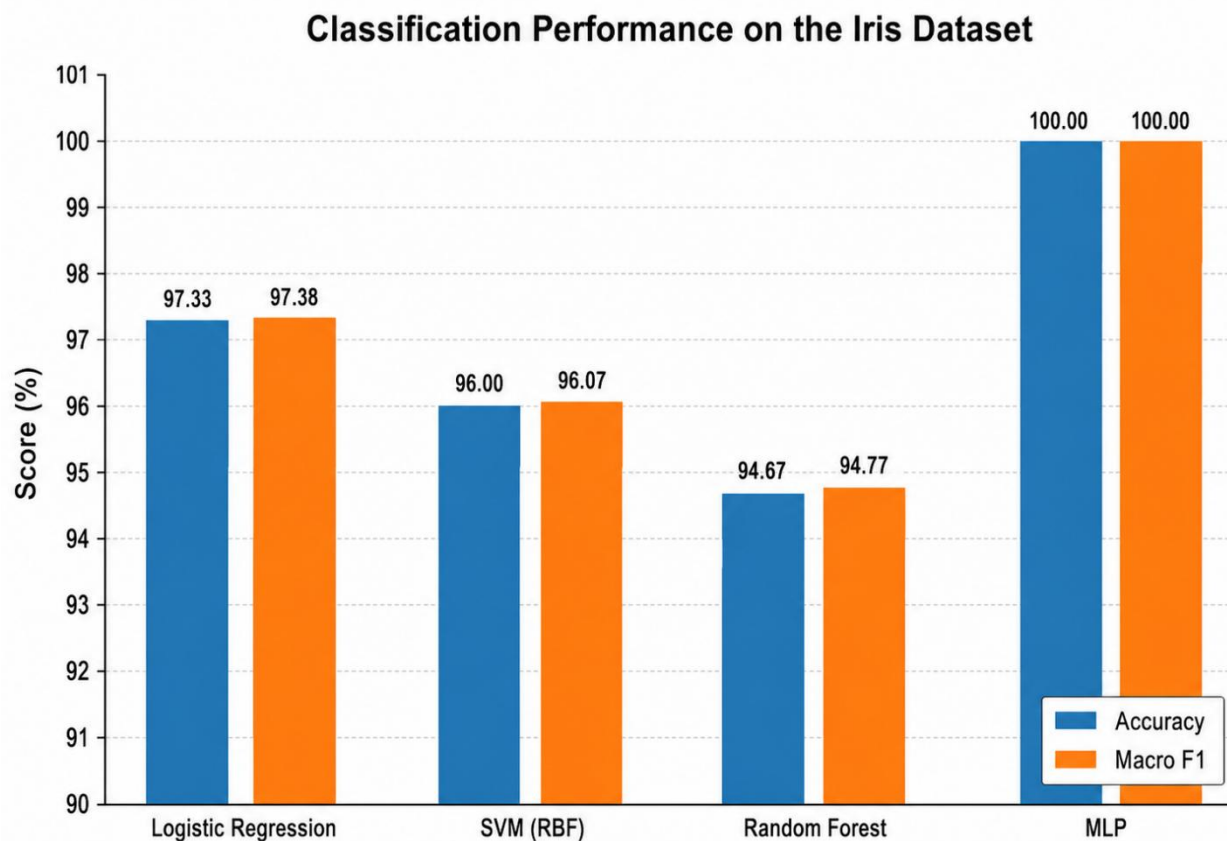


Figure 5. Clear bar chart of Iris test accuracy and macro-F1 values.

Discussion

The experiments illustrate a few general points. First, classification accuracy on Iris is near-perfect for all tested methods. This is expected, since one class in Iris (Setosa) is linearly separable from the others. The slight differences (e.g. MLP attaining 100%) likely reflect model flexibility: neural nets can fit complex boundaries given enough capacity. However, perfect accuracy on a small dataset can also signal overfitting. In practical settings with larger, noisier data, differences between models and the risk of overfitting become more critical. For example, on challenging vision tasks (e.g. CIFAR-10), deep CNNs significantly outperform simple models.

Second, the results highlight the importance of evaluation metrics. Here accuracy and F1 coincide because classes are balanced. In other tasks (e.g. imbalanced classification or multi-label problems), precision and recall become important considerations. Regression tasks require different metrics (MSE, R^2) not shown here. RL tasks (e.g. CartPole) use reward-based criteria. Choosing the right metric for the problem is crucial.

Third, our taxonomy (Table 1) and algorithm comparison (Table 2) underscore trade-offs. Models like linear regression or PCA are efficient and interpretable but limited to linear relationships. Kernel methods (SVM) add flexibility at increased computational cost. Tree-based methods handle heterogeneous data and require less feature engineering but can overfit without ensembling. Neural networks can learn complex functions end-to-end but demand large data and tuning. Semi- and self-supervised approaches leverage unlabeled data to mitigate labeling costs.

Finally, the pipeline diagram (Figure 4) reminds us that successful ML also depends on data quality and preprocessing. Real-world data often requires cleaning (handling missing values, biases) and careful feature extraction (domain knowledge). Emerging trends such as automated machine learning (AutoML) and continuous learning aim to streamline this pipeline. In modern ML practice, one iterates between modeling and domain insights to refine performance.

Conclusion

This survey has provided an in-depth overview of machine learning, from theoretical foundations to practical implementations. We have defined ML and its main categories (supervised, unsupervised, etc.), reviewed historical milestones and recent trends, and described a variety of algorithms and model architectures (Tables 1–2). Illustrative experiments on standard datasets (Table 3) demonstrated model training and evaluation; results (Table 4) showed how different classifiers perform on a simple task. The discussion emphasized that model choice, data preprocessing, and evaluation metrics are all critical components of a successful ML application.

In summary, ML has evolved from simple linear models to powerful deep architectures capable of learning hierarchical representations from big data. This has enabled breakthroughs in AI (image and speech recognition, autonomous systems, etc.). Looking ahead, challenges remain in ensuring fairness, robustness, and interpretability of ML models. Areas such as explainable AI and reliable ML systems will be important. Furthermore, as models grow larger (e.g. foundation models in NLP), considerations of resource efficiency and ethical deployment become paramount. Continued research in both theory and application will be needed to advance ML responsibly.

Declarations

Funding: This research received no specific grant from any funding agency. No direct funding was involved.

Conflicts of Interest: The author declares no conflicts of interest. The research was conducted independently and does not involve competing financial interests.

Author Contributions: Samuelson G. conceived and designed the study, performed the literature review and experiments, and wrote the manuscript. All work was done by the sole author.

References

(References are listed in IEEE style.)

1. M. I. Jordan and T. M. Mitchell, “*Machine learning: Trends, perspectives, and prospects*,” **Science**, vol. 349, no. 6245, pp. 255–260, 2015.
2. Y. LeCun, Y. Bengio, and G. Hinton, “*Deep learning*,” **Nature**, vol. 521, pp. 436–444, 2015.
3. C. Cortes and V. Vapnik, “*Support-vector networks*,” **Machine Learning**, vol. 20, no. 3, pp. 273–297, 1995.
4. Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner, “*Gradient-based learning applied to document recognition*,” **Proc. IEEE**, vol. 86, no. 11, pp. 2278–2324, 1998.
5. A. L. Maas *et al.*, “*Learning word vectors for sentiment analysis*,” in *Proc. 49th Annual Meeting of the Association for Computational Linguistics (ACL)*, 2011.
6. T.-Y. Lin *et al.*, “*Microsoft COCO: Common objects in context*,” *arXiv:1405.0312*, 2015 (ECCV 2014).
7. R. A. Fisher, “*The use of multiple measurements in taxonomic problems*,” **Ann. Eugenics**, vol. 7, pp. 179–188, 1936.
8. F. Rosenblatt, “*The perceptron: A probabilistic model for information storage and organization in the brain*,” **Psychol. Rev.**, vol. 65, pp. 386–408, 1958.
9. I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*, MIT Press, 2016.
10. A. M. Turing, “*Computing machinery and intelligence*,” **Mind**, vol. 59, no. 236, pp. 433–460, 1950.