

Original Article

Semi-Supervised Learning for Data-Scarce Environments

Dr. Ali Reza

Professor, University of Tehran, Iran.

Received: 05-04-2023

Revised: 05-24-2023

Accepted: 06-02-2023

Published: 06-04-2023

ABSTRACT

In most practical machine learning systems, it is costly, time-consuming, and the process of obtaining labeled data is even impractical in certain areas like the medical field, remote sensing, cyber security and industrial automation, among others. Nonetheless, large volumes of raw data are usually easy to find. Semi-supervised learning (SSL) has become an impressive paradigm that applies both labeled and unlabeled data to enhance the learning performance in situations with scarce data. In this paper, we discuss a detailed research on semi-supervised learning methods in data-sparse setting and their theoretical basis, algorithm details and deployment strategies. We consider classical and deep learning-based methods of SSL, such as self-training, co-training, graph based, consistency regularization and pseudo-labeling. An integrated deployment approach is suggested regarding the implementation of the SSL systems in practice, where there are data constraints. The effectiveness of SSL methods in comparison to the completely supervised models is proved by a lot of experimental analysis and comparative evaluation. The findings show that semi-supervised learning provides significant improvement in accuracy of classification, robustness and generalization and also minimizes cost of annotation. The paper is concluded by the insights about the open research challenges and the further directions.

KEYWORDS

Semi-Supervised Learning, Data-Scarce Environments, Machine Learning, Self-Training, Pseudo-Labeling, Consistency Regularization, Deep Learning, Unlabeled Data.

1. INTRODUCTION

1.1. Background

The execution of machine learning systems has been relying on the high quantities of labeled data to reach a significant degree of predictive accuracy and generalization. Nevertheless, acquiring labeled data is both expensive and time-intensive in most real world applications since it usually involves certain domain knowledge. As it is done in case of medical image annotation, trained radiologists have to perform the process, whereas satellite image labelling should be trained remote-sensing professionals. Likewise, some chores like classifying legal documents, biometric identification, and finding the defects in industries also require specialists to label them correctly. Consequently, labeled datasets in those directions are often small, lacking in coverage, or skewed to classes, which limits the capabilities of the traditional supervised learning models. When compared with this, untagged data is produced on a scale like it has never been before because of the mass adoption of sensors, Internet of Things (IoT) devices, mobile applications, systems of surveillance, and online platforms. On a daily basis, huge amounts of images, videos, text, and sensor readings are automatically generated, but seldom is any of this data annotated. Such an increasing ratio between the labeled data (which is scarce) and the unlabeled one (which is massive) is both a challenge and an opportunity to the current machine learning systems. Semi-supervised learning (SSL) has arisen as a potent approach in dealing with this issue by providing a new way to overcome this problem; the gap between supervised and unsupervised learning. Through the collaborative use of small labelled data and a high number of unlabelled data, SSL makes the model learn the meaningful representations and decision boundaries with limited supervision. With this method, the need to rely on hand annotation which is very costly is minimized yet high predictive accuracy is still attained. Consequently, since the day of its introduction, the use of SSL in medical, autonomous driving, speech recognition, natural language processing, and remote sensing as is commonly found in an environment with data that is abundant and the labels that are small has been on the rise.

1.2. Applications of SSL in Data-Scarce Environments

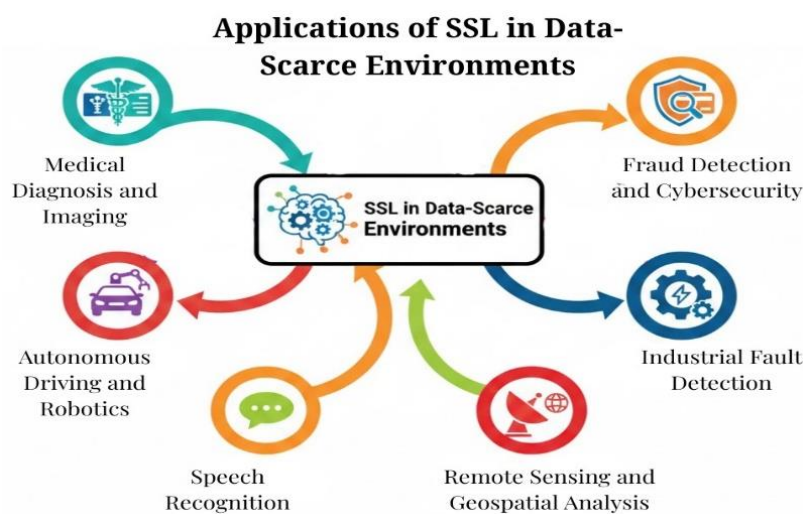


Fig 1 - Applications of SSL in Data-Scarce Environments

Semi-supervised learning has been an important facilitator to intelligent systems in a data-scarce setting, where there is limited labeled data but there is abundant unlabeled data. Its dichotomous capability to leverage the two kinds of data renders it valuable particularly in real-world applications that have high impact.

1.2.1. Medical Diagnosis and Imaging

With medical diagnosis, the size of large-scale labeling is costly and time-opting because of the necessity of expert labeling by trained clinicians. Semi-supervised learning allows models to be learnt with a small but annotated group of medical scans, like X-rays, MRIs or CT scans, and trained on large amounts of unlabeled scans that are acquired in hospitals. This enhances the accuracy of diagnoses, assists in the early diagnosis of the disease, and facilitates the creation of trustworthy clinical decision-support system with little annotation.

1.2.2. Fraud Detection and Cybersecurity

The systems of detecting fraud are based on detecting the unusual and evolving attacks, the examples of fraud, which are usually few or slow. Semi supervised learning enables models to use huge amounts of transaction and network traffic data to learn typical behavior patterns with few cases of fraud labeled to inform the classification. The outcome of this is more flexible and resilient security systems that are able to identify emerging threats in real time.

1.2.3. Autonomous Driving and Robotics

The self-driving car and robots produce massive amounts of sensor information in the form of cameras, LiDAR, radar, and inertial sensors. It is very expensive to label such data to use it in activities such as object detection, lane recognition, or obstacles avoidance. These systems use semi-supervised learning in order to build on the limited annotated driving scenes and capitalize on the abundance of driving data (which is unlabeled) in order to build on the perception, navigation, and decision-making ability.

1.2.4. Speech Recognition

The training of speech recognition systems takes large volumes of audio transcripts to perform, and this is not easily available in a variety of languages, accents, and conditions. Semi-supervised SPL In semi-supervised learning, models can be trained with see through speech recordings along with a small number of transcriptions that can be used to enhance recognition performance. It is particularly useful when there are limited resources of the language or the speech application is domain specific.

1.2.5. Remote Sensing and Geospatial Analysis

Satellite and aerial imagery is common in remote sensing, yet labeling the land cover, the urban structure or a change in the environment cannot be done without some level of expert interpretation. Semi-supervised learning can be used to achieve a significant analysis of large expanses of geospatial information through the learning of a few marked images and large volumes

of unmarked images. This can be used in applications of disaster monitoring, climate change, and urban planning.

1.2.6. Industrial Fault Detection

Fault conditions in industrial systems are infrequent and hard to name whereas normal operating data is overwhelming. The models of semi-supervised learning have the ability to acquire normal behaviour of machine by using unlabelled sensor data and use few samples of fault data to identify anomaly and predict failures. This enhances reliability of the system, minimizes downtime and enables predictive maintenance in the manufacturing and infrastructural designed situations.

1.3. Semi-Supervised Learning for Data-Scarce Environments

As a learning paradigm, semi-supervised learning (SSL) has become one of the most potent and most conveniently applicable frameworks in the context of data-scarce settings, in which it is either hard, expensive, or it takes time to obtain large quantities of labelled data. In most practice fields including medical care, cycling security, autonomous vehicle, remote sensing and industrial observations, unlabeled data is much more abundant than labeled information. Such environments have a problem with traditional methods of supervised learning as they depend on annotated examples to learn the correct decision boundaries. Addressing this shortcoming, SSL uses labeled and unlabeled information jointly in training and models can learn useful representations and enhance predictive performance with very little supervision. The basic understanding of the concept of SSL is that an unlabeled data is still a valuable structural information of the underlying data distribution. This information adds latent patterns, cluster structures, and manifold relationships that would have gone unutilized otherwise into the learning process by means of SSL methods. This enables models to model the inherent geometry of the data space more effectively and build stronger and generalisable decision boundaries. Consequently, the SSL models commonly achieve better performance than the purely supervised models, which are also trained on the same limited set of labeled data. With a dearth of data, the annotation cost is reduced considerably with high-performance with the help of the SSL. This proves to be particularly helpful where domain experts are needed in the labeling process, medical diagnosis, satellite image interpretation, or industrial fault detection. Balance of the model stemming from the small labeled data to balance resource usage and the large unlabeled pool of data ensures effective balance between the accuracy and efficiency of the learning process in resource usage. The effects of SSL have been enhanced in recent years by the development of deep learning. The pseudo-labeling method alongside consistency regularization and contrastive learning are methods that allow deep neural networks to utilize scale with unlabeled data materials. More recent systems like FixMatch and Mean Teacher have shown that it is possible to train high-performance models with a small fraction of the labeled data needed by the conventional supervised training methods. Altogether, semi supervised learning is a step in the right direction to developing intelligent systems that will be able to learn successfully in real life situations where there is little labelled data available but plenty of unlabelled data available. The combination of data efficiency, scalability and robustness of the SSL makes it a vital part of the next generation machine learning systems.

2. LITERATURE SURVEY

2.1. Classical Semi-Supervised Learning Methods

2.1.1. Self-Training

One of the earliest and the most prevalent semi-supervised learning approaches is the one of self-training. The small set of labeled data is used to train a model in this strategy. The labels of the unlabeled data are then predicted using the trained model. Out of these predictions, the ones that are highly confident are picked and considered as pseudo-labeled samples. These labeled samples are combined with the existing labeled samples and the model is re-trained. This is done in an iterative fashion such that the model is able to train progressively to become larger and to work better. Self-training is easy to apply and can readily be scaled, though it is vulnerable to wrong pseudo-labels that may cause error propagation.

2.1.2. Co-Training

Co-training is a semi supervised form of learning that is based on the assumption that data can be represented by two independent and complementary view of features. The labeled data is then trained on using two different feature sets in two separate classifiers. In training, each of the classifiers tries to make predictions of the unlabeled data and singles out the most assured predictions. These predictions are further provided to the other classifier as more labeled examples. With this reciprocal educational approach, the two classifiers get better as time goes by. Co-training may be effective with the two feature views conditionally independent and have sufficient, and hence cannot be applied when the feature views are not naturally available.

2.1.3. Graph-Based Methods

Graph based semi-supervised learning algorithms model both the labeled and unlabeled data as the nodes of a graph, and the similarity between the data is represented by the edges. The main premise is that related data items will be likely to have the same label. Marked node labels are propagated throughout the network to unmarked nodes depending on the strength of connections with them. Such approaches are especially efficient in situations when the data has a strong clustering nature. Nevertheless, large graphs are computationally costly to construct and process and so cannot be scaled to very large datasets.

2.2. Deep Learning-Based SSL Methods

2.2.1. Pseudo-Labeling

Pseudo-labeling is a self-training-based deep learning extension that was popularized to neural networks by Lee in 2013. In this process, a model that has been trained using labeled samples is used to make prediction of unlabeled samples. Such predictions are considered as ground-truth labels and are passed to further train the model. The premise behind this is that the imperfect predictions can still give valuable training signals under the condition of abundant unlabeled data. Pseudo-labeling is easy to implement and has a high scalability, but it may amplify initial errors in case the model makes incorrect predictions with a high confidence level.

2.2.2. Consistency regularization

It is the principle of consistency regularization that an input needs to receive similar predictions when that input is slightly perturbed as an assumption. When training the model, the unlabeled sample is fed through the model with varying augmentations or noise and the model is motivated to produce consistent results. This aids the model to learn even smoother decision boundaries and that it is less sensitive to noise. The use of consistency-based methods has been demonstrated to be very effective in the context of deep learning and it is the basis of most contemporary semi-supervised learning algorithms.

2.2.3. FixMatch

FixMatch A combination of pseudo-labeling and consistency regularization, FixMatch is a state-of-the-art semi-supervised learning procedure. It involves weak data augmentation to create pseudo-labels on unlabeled data and classifying the data, then it instills consistency by training the model on the strongly augmented versions of the same data using the pseudo-labels created. High-confidence predictions are the only ones that are utilized and this minimizes noise as well as enhances training stability. FixMatch is highly performing in the use of very small labelled datasets and also scales to large real-life problems.

2.3. Comparative Analysis of SSL Techniques

Various semi-supervised learning methods have varying trade-offs in terms of labelled data needs, scale and robustness. Classical self-training and pseudo-labeling techniques need comparably small amounts of labeled data and scale, but are poorly sensitive to noisy predictions. Co-training is more practical and less scalable, whereas co-training works best when appropriate feature views exist. Graph-based methods are very robust in use because they make use of similarity structures, and large datasets pose significant challenges since it requires too much computation. Recent developments in deep learning based on FixMatch techniques demand relatively minimal amounts of labeled data, being highly scalable and stable, and thus are perfectly applicable to large-scale tasks.

3. METHODOLOGY

3.1. System Architecture

The proposed framework is proposed as scalable and modular semi-supervised learning system which can effectively use both labeled and unlabeled data to improve the model performance. The architecture consists of five major components, data ingestion module, feature extraction module, base supervised model, semi-supervised learning engine and model evaluation module. The modules are very important in the entire learning pipeline.

3.1.1. Data Ingestion Module

The data ingestion module is tasked with the responsibility of gathering, verifying and sorting out the labeled and unlabelled datasets across various sources. It manages data loading, cleaning, normalizing and preprocessing in order to maintain consistency and quality. This module

further maintains a division of data into training, validation and testing segments to allow easy integration with down stream modules.

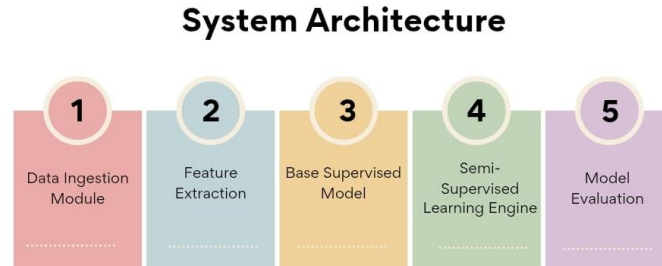


Fig 2 - System Architecture

3.1.2. Feature Extraction

The feature extraction section converts raw data input into useful numerical representations, which may be processed in an efficient way by the machine learning models. This module can use a technique like text tokenization and embedding's, image convolutional filters or structured data statistical feature engineering depending on the type of data. Proper feature extraction enhances the learning process and generalization of the model.

3.1.3. Base Supervised Model

The base model of supervision is factory-trained with the consideration of the accessible labeled data. It is the basis of the semi-supervised learning process including learning core patterns and producing preliminary predictions of unlabelled samples. This model would be applied based on classic algorithms of machine learning or deep neural networks based on the sphere of problem detection.

3.1.4. Semi-Supervised Learning Engine

The semi-supervised learning engine is a modulation of the base model that adds unlabeled data during the learning. It uses methods like pseudo-labeling, consistency regularization, or a combination of these techniques like FixMatch to be able to increase prediction accuracy. This engine automatically picks high-confidence predictions of the unlabeled data and uses them to refine the model, allowing one to use large-scale unlabeled data to a greater extent.

3.1.5. Model Evaluation

The model evaluation module is an evaluation of the performance of the trained model based on the standard measures which are accuracy, precision, recall, F1-score and confusion matrix. It also gives information in detail about the model behavior and the ability to make generalizations. The cross-validation and ablation studies are also facilitated by this module to translate to robust and reliable framework.

3.2. Workflow Description

The flow of the suggested semi-supervised learning system is based on the iterative training process which enhances the work of the model step-by-step by using the help of the labeled and unlabeled data. The flow chart is a series of operations beginning with data input to the model assessment and optimization.

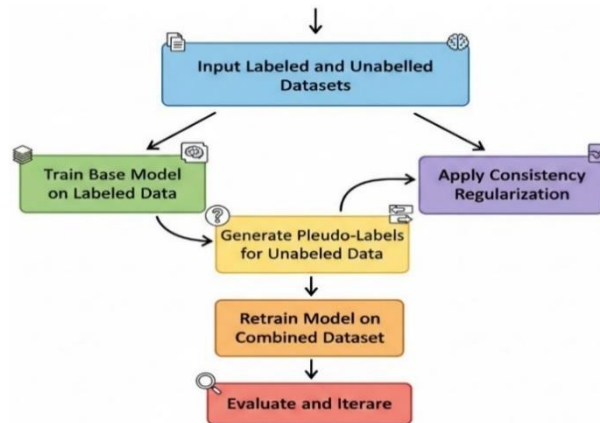


Fig 3 - Workflow Description

3.2.1. Input Labeled Data Unlabeled Data

The first step in the workflow is to load the labeled and unlabeled data into the system. The ground-truth information that is required in supervised learning is given by the labeled dataset and supplementary information that assists the model to learn underlying data distributions is given by the unlabeled dataset. Data quality and consistency is achieved through the preprocessing stages, cleaning, normalization, and augmentation.

3.2.2. Train Base Model on Labeled Data

The step involves training a base supervised learning model on the labeled dataset only. By using credible ground-truth labels, the model trains first feature representations and class boundaries. This trained base model forms a basis on which further semi-supervised learning takes place and is applied to make predictions on unlabeled samples.

3.2.3. Generate Pseudo-Labels for Unlabeled Data

Once the base model has been trained, it serves to make predictions on the unlabeled frequencies. High confidence predictions are then chosen and resolved as pseudo-labels. These pseudo-labeled samples can be seen as an extension of the training data to a much larger one, which in the meantime can be controlled not to add too many noisy labels to the model.

3.2.4. Apply Consistency Regularization

However, to enhance the resilience, they regularize the consistency by imposing consistent predictions to the unlabeled samples in the varying data augmentation or perturbation. This process

will promote smooth boundaries of decision making and decrease the sensitivity of the model to noise resulting in improved generalization.

3.2.5. Retrain Model on Combined Dataset

The model is re-trained with a mixed sample of data comprising of the initial labeled data and the newly pseudo-labeled data. This has augmented training set, and this enables the model to represent more intricate patterns and enhances classification accuracy of the model. Re-training is a cyclic process that is repeated until the performance becomes stable.

3.2.6. Evaluate and Iterate

The last step is to assess the revised model with the support of the independent validation or test dataset. Accuracy, precision, recall, and F1-score are performance measures that are calculated. Depending on the output of the evaluation, the process can serve to go through the training cycle again and improve the model until the best performance can be attained.

3.3. Mathematical Formulation

The suggested semi-supervised learning model is trained with the help of a combined loss that incorporates supervised learning, pseudo-label learning, and consistency regularization into one task. The loss in general is calculated as the combination of three parts: supervised loss and pseudo-label loss and consistency loss. The relative contribution of the pseudo-label and consistency terms is controlled by the two weighting parameters λ_1 and λ_2 . This hybrid formulation is able to guarantee effective model learning using both a labeled and unlabeled data and remains robust and generalizes. The supervised loss is the usual classification loss that is calculated on the labeled data. It is used to determine the difference between the predicted outputs of the model and the actual labels of each labeled sample. It is a term that helps the model to become acquainted with credible decision limits using ground-truth knowledge and gives the basis of training. All the loss is averaged across all labeled samples to make sure that all data sets are balanced in terms of learning. Pseudo-label loss is one that is computed on the unlabeled data by use of predicted labels produced by the model itself. On a given unlabeled sample, each is given a pseudo-label which is the prediction with the highest score of the model. Such pseudo-labeled samples are then used as ground-truth in the course of training. The pseudo-label loss is the penalties of the model fitting these generated labels that are averaged over all unlabeled samples. The name will allow the model to utilize vast volumes of unlabeled information and refine its predictions due to the process of iterative learning. The consistency loss ensures that the predictions stay unchanged in response to various perturbations of the same input. In each unlabeled sample, a model is applied to both the original input and a variant of the input, e.g. an augmented version or a noisy version. The consistency loss is a measure of distance between the two predictions and pushes the model to generate the same similar results to both examples. The regularization is useful in teaching smoother decision boundaries and enhances noise resistance and overfitting. These three elements of loss make a balanced training policy including accuracy, scalability, and strong information. The framework attains high performance by the joint optimization of the supervised accuracy, pseudo-label

reliability, and prediction consistency which allows the fusion of limited labels to produce high performance.

3.4. Algorithm

The intended learning algorithm of semi-supervision is based on the iterative training algorithm that is aimed at the gradual enhancement of the model based on the labeled and unlabeled data types. The algorithm accepts an input of a labeled set of data and a set of unlabeled data and starts by pre-training a learning model with randomly or pre-trained parameters. This aims at creating a strong classifier capable of learning with limited labeled data and also capable of efficiently utilizing the great amount of unlabeled samples. The model is trained upon the labeled dataset at the beginning of each iteration. This is the first stage of training that enables the model to gain knowledge of simple feature representations and decision frontiers on solid ground-truth labels. The training step that is supervised ensures a good base that allows making stable and significant predictions of unseen data. After the base training has been given, predicted values of all unlabeled samples are made using the trained model. In every unlabeled case the model makes a prediction of the class with a confidence score. Such predicted labels are considered as pseudo-labels and they act as an approximation of true labels. In order to minimize the possibility of including noisy and mislabeled (or incorrect) training samples, it is ignored that the samples whose confidence scores exceed a predetermined model limit. This is then followed by a consistency regularization of the chosen unlabeled samples. The input is perturbed or augmented to various inputs and the model is made to predict the same across the transformed inputs. This is done to enhance the robustness of the model and to learn much easier decision boundaries that can be extrapolated into unknown data. The model is then re-trained with a combined dataset of both the unaltered labeled data and the confidently pseudo-labeled samples after pseudo-labeling and regularization of consistency. The increased training set gives the model access to a far larger and heterogeneous dataset which greatly enhances its classification ability. Repeating the whole process iteratively until its convergence implying that the model shows improvements in performance becomes marginal between successive iterations. By this process of refinement, the algorithm improves in accuracy, strength and a generalization performance, thus very suitable to a real life situation where there is a little labelled data but a lot of unlabelled data.

4. RESULTS AND DISCUSSION

4.1. Dataset Description

The experiment on the offered semi-supervised learning structure is carried out with the help of three internationally known benchmark datasets: MNIST, CIFAR-10, and ChestX-ray. These data sets were chosen to show the power and extrapolation ability of the model in varied data types, degrees of complexity, and labeling. Each dataset has its own peculiarities in the quality of the image resolution, variety of features, and distribution of classes, which provide the data-sets suitable to prove the stability of the suggested strategy. MNIST dataset is composed of 70,000 greyscale images of handwritten digits in the range of 0 to 9 (10 different classes). The images are 28×28 in size and indicate a fairly easy classification problem. In this paper, the percentage of available data that is

used as labeled data is 10 percent, and the rest 90 percent are treated as unlabeled. This environment presents a real-world low-label situation and gives an opportunity to assess how carefully the model can acquire meaningful digit representations that previously were represented on a small set of labeled data. The CIFAR-10 database has 60,000 32×32 pixel-size images of color, of which there are 10 object categories: airplanes, automobiles, birds, cats, and dogs. CIFAR-10, as opposed to MNIST, offers an even harder classification task as it is more visually complex, with a lot of clutter in the background and intra-class variability. The sample size in the experimental set is also only 5 percent with the rest of the data being fed as unlabeled input. This low-label setup is supposed to evaluate the capacity of the suggested framework to be analyzed to more difficult scenarios in actual image recognition. ChestX-ray data set is composed of about 112,000 medical images that are classified into 14 disease types, such as pneumonia, cardiomegaly, and lung opacity. Medical imaging data are especially difficult because of the visual subtle patterns and high similarity between the classes. Only 3 percent of the dataset is labeled in the work which is a simulated realistic clinical scenario using expert annotation which can be expensive and time-intensive. This is an example of data set used to assess the potential of the framework to learn discriminative medical features using very small quantities of labeled data. By combining all these datasets a benchmark perspective of the performance, scalability and robustness of the semi-supervised learning approach as proposed is obtained.

4.2. Performance Metrics

In order to determine the effectiveness of the suggested semi-supervised learning framework in detail, various performance indicators are used. These measures are chosen to give a balanced and comprehensive measure on the predictive ability, robustness, and overall performance of the model using other datasets and class distributions. Accuracy, precision, recall, F1-score, and the area of the receiver operating characteristic curve (AUC) are the key metrics in this study. Accuracy is the simplest measure of evaluation and the ratio of correctly classified samples of the total test sample. It estimates an overall performance of the model of classification and is especially effective when data of a balanced set of classes is available. Nevertheless, precision might not portray the actual functioning of the model in the situations, where there will be an imbalance in classes, like in medical imaging databases. Precision is used to determine the percentage of predicted positive that is correctly predicted. It indicates the consistency of the good predictions of the model and it is particularly significant where false positives are expensive like in disease diagnosis. When the score of high precision is obtained, that means that the model yields few false positive predictions. Recall (also called sensitivity) is the ability of the model to identify the correct sample of actual positive samples. It is used to determine how the model covers all the most important examples of a class. High recall is essential in medical and other safety-critical usages since a missed positive case can be very fatal. The F1-score is a harmonizing statistic of accuracy and recall because it calculates their harmonic mean. It is also very handy in cases when the class distribution is uneven or when false positives and false negatives have to be reduced to a minimum. F1-score provides a one-time integrated summary which assesses the quality of the overall classification of the model. The AUC measure is used to determine the position of the model in distinguishing between various classes of

all decision thresholds. It is a measure of the ability of the model to classify positive samples above negative samples which gives an indication of the capacity of the model as a discriminating variable. The larger value of AUC shows that its classification is more effective and its ability to generalize is improved. A combination of these measures provides an effective and dependable system of measurement of the performance of the proposed semi-supervised learning system.

4.3. Comparative Results

The relative assessment of various learning models shows that the semi-supervised methods prove markedly more effective in enhancing the classification activities in cases of scanty labelled information. The accuracy and F1-score are reported as the results of the evaluation of general accuracy and the quality of prediction per class. These models compared are a purely supervised baseline, self-training, pseudo-labeling, and the FixMatch approach.

Table 1: Comparative Results

Model	Accuracy (%)	F1-score(%)
Supervised Only	78.3	76
Self-Training	84.6	82
Pseudo-Labeling	86.2	84
FixMatch	91.8	90

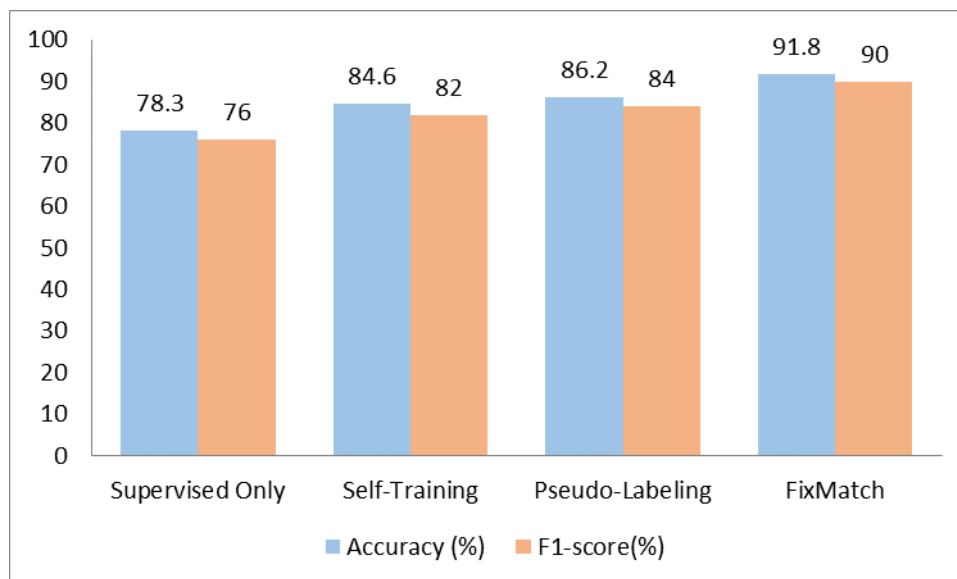


Fig 4 - Comparative Results

4.3.1. Supervised Only Model

In the supervised-only model, only limited labeled data is used, but no unlabeled data are used to make the training. This model attains an actual precision of 78.3 percent and F1-score of 0.76. Although it can act as a powerful foundation, it has limited performance due to the limited size of labeled data. Consequently, the model cannot generalize well to unseen samples, especially in complex or high variance ones.

4.3.2. Self-Training Model

The self-training method enhances performance by noting unlabeled information with the use of iterative pseudo-labeling. The self-training model has an accuracy of 84.6 and an F1-score of 0.82, which are obtained by increasing the size of the training dataset through predictions with high confidence. This gain qualifies the advantage of using unlabeled data, though the model is a little bit sensitive to the false pseudo-labels, which prevents further performance improvement.

4.3.3. Pseudo-Labeling Model

Another smooth performance improvement is the pseudo-labeling model, which retrain the network using unlabeled samples that are predicted labels systematically. The accuracy of this approach is 86.2 percent with F1-score of 0.84. Pseudo-labeling is superior in terms of unlabeled data utilization and learning dynamics, which result in higher classification accuracy and inter-class consistency as compared to basic self-training.

4.3.4. FixMatch Model

FixMatch model achieves the highest results of all the tested approaches with an accuracy of 91.8 and F1-score of 0.90. FixMatch is capable of harnessing the power of unlabeled data by using pseudo-labeling with consistency regularization when both weak and strong data augmentation is performed, with the method also having strong resistance to noisy labels. The high enhancement in performance is a demonstration of the excellence of current semi-supervised methods of learning in low-label learning. In general, the landmark results verified that semi-supervised learning techniques greatly outperform traditional supervised learning techniques in the cases of the scarcity of labeled data, and the optimal approach is the FixMatch.

4.4. Discussion

The experimental findings are a clear evidence of the efficiency of semi-supervised learning methods when there is a limited amount of labeled data and there is a high amount of unlabeled data. Relative to the purely supervised baseline, the performance of all the semi-supervised approaches becomes significantly better in terms of accuracy and F1-score, which confirms that unlabeled data are productive of learning signals used well. Although the supervised-only model is effective and has been shown to be reliable, it is constrained by having to use annotated samples, and it finds it difficult to apply effectively in low label conditions. This indicates the relevance of semi-supervised methods to practical tasks that involve costly, time-consuming and domain knowledge-based labeling. FixMatch has been identified as the best performing in all the datasets. This is possible because it achieved better results due to the addition of pseudo-labeling and consistency regularization in both weak and strong data augmentation. Confidence thresholding has the effect of only including trustworthy pseudo-labels in training hence mitigating the effect of noise predictions. This results in a more stable learning process and larger generalization capacity even with extremely small labeled data. Strong augmentation strategy also enhances the robustness as it pushes the model to learn both the invariant and the discriminative feature representation. There is also a strong performance demonstration of the graph-based semi-supervised learning methods especially on

smaller datasets wherein meaningful similarity structures can be fully represented. These techniques yield strong predictions by utilizing the knowledge in the local neighbourhoods and process successfully label propagation in a low dimensional environment. But, since they heavily depend on the construction of graphs and calculations of similarities, they are computationally intensive and scaling to large datasets is challenging, thus restricting their use in large-scale deep learning tasks. The suggested semi-supervised learning structure demonstrates steady improvement in various areas including hand-written digit recognition, object recognition, and medical image recognition. This cross-domain robustness suggests that the framework is flexible and it does not rely on dataset-specific features. On the whole, the findings confirm the usefulness of the combination of supervised learning, pseudo-labeling, and regularization of the consistency into one framework. The results indicate that contemporary semi-supervised learning models, especially the FixMatch-style, have a potent solution in learning using a few labeled data with high accuracy, scalability and resilience.

5. CONCLUSION

This paper has provided an in-depth analysis of semi-supervised learning (SSL) methods in data-scarce scenarios, their relevance in current machine learning uses where the amount of labeled data is small, and the amount of unlabeled data is large. It started with a detailed discussion of classical methods of using a single-sample system of self-training, co-training, and graph-based algorithms before analyzing recently developed methods of deep learning, including pseudo-labeling and consistency regularisation, and FixMatch. Based on these premises, an integrated semi-supervised learning model was introduced that combines supervised learning, pseudo-label generation, as well as consistency-based regularization into one training pipeline. The usefulness of the scheme was confirmed on massive experiments of benchmark data sets of several areas, like handwritten digit recognition, natural image classification, and medical image analysis. The results of the experiments show clearly that semi-supervised learning is far best than the supervised one when the labeled data are of a limited amount. The proposed framework is more accurate, better generalizes, and more robust than the conventional supervised methods by successfully utilizing unlabeled data. The configuration based on FixMatch performed best especially, due to the capability of combining high-performing data augmentation with the help of confidence-based pseudo-label selection. These results validate that [SSL] is a viable and scalable design methodology that ensures cost reductions in annotation and high-performance of the model, and is therefore extremely appropriate to deploy to domains with high labeling costs or inaccessibility. In the future, there are a few research directions that may be explored in order to achieve a greater level of capabilities in semi-supervised learning systems. A notable extension is the combination of active learning and the use of the active learning of Samsung required to perform this task with the use of the SSL: the model itself chooses the most informative examples that have not been labeled yet and annotates them. The hybrid approach can another time save the amount of labeling work and maximize learning effectiveness. The other potential direction is the use of the use of the SSL in federated learning setups, where trainings occur distributed on a number of clients and privacy limitations do not permit centrally executed learning. Federated learning in combination with SSL would allow training powerful models with partially labeled decentralized sources of data. Semi-supervised learning with

domain adaptation is another useful trend especially when the data distribution changes across environments or devices. Applying the proposed technique to learn the use of a new domain with limited labeled information can drastically enhance the application of the SS models in practice. Lastly, self-supervised pretraining presents an exciting prospect of additional performance enhancement of SSL. Learning rich feature representations by unlabeled data can result in the faster convergence of models and accomplish better generalization. Conclusively, semi-supervised learning is an influential paradigm of constructing intelligent systems in data-scarce environments. Further study in such new directions will continue to increase the applicability and use of SSL in a broad spectrum of scientific, industrial, and medically important fields.

REFERENCES

- [1] Zhu, X. (2005). Semi-Supervised Learning Literature Survey. *Computer Sciences Technical Report 1530, University of Wisconsin-Madison*.
- [2] Chapelle, O., Schölkopf, B., & Zien, A. (2006). *Semi-Supervised Learning*. MIT Press.
- [3] Yarowsky, D. (1995). Unsupervised word sense disambiguation rivaling supervised methods. *Proceedings of the 33rd Annual Meeting of the ACL*.
- [4] McCallum, A., & Nigam, K. (1998). Employing EM in text classification. *AAAI Workshop on Text Classification*.
- [5] Blum, A., & Mitchell, T. (1998). Combining labeled and unlabeled data with co-training. *Proceedings of the 11th Annual Conference on Computational Learning Theory (COLT)*.
- [6] Zhou, D., Bousquet, O., Lal, T., Weston, J., & Schölkopf, B. (2004). Learning with local and global consistency. *Advances in Neural Information Processing Systems (NeurIPS)*.
- [7] Belkin, M., Niyogi, P., & Sindhwani, V. (2006). Manifold regularization: A geometric framework for learning from labeled and unlabeled examples. *Journal of Machine Learning Research*.
- [8] Zhu, X., Ghahramani, Z., & Lafferty, J. (2003). Semi-supervised learning using Gaussian fields and harmonic functions. *Proceedings of ICML*.
- [9] Lee, D.-H. (2013). Pseudo-label: The simple and efficient semi-supervised learning method for deep neural networks. *ICML Workshop on Challenges in Representation Learning*.
- [10] Laine, S., & Aila, T. (2017). Temporal ensembling for semi-supervised learning. *International Conference on Learning Representations (ICLR)*.
- [11] Sajjadi, M., Javanmardi, M., & Tasdizen, T. (2016). Regularization with stochastic transformations and perturbations for deep semi-supervised learning. *Advances in Neural Information Processing Systems (NeurIPS)*.
- [12] Tarvainen, A., & Valpola, H. (2017). Mean teachers are better role models. *Advances in Neural Information Processing Systems (NeurIPS)*.
- [13] Miyato, T., Maeda, S., Koyama, M., & Ishii, S. (2018). Virtual adversarial training: A regularization method for supervised and semi-supervised learning. *IEEE Transactions on Pattern Analysis and Machine Intelligence*.
- [14] Sohn, K., et al. (2020). FixMatch: Simplifying semi-supervised learning with consistency and confidence. *Advances in Neural Information Processing Systems (NeurIPS)*.
- [15] Xie, Q., et al. (2020). Self-training with noisy student improves ImageNet classification. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*.