

Machine Learning

| Semi-supervised and Self-supervised Learning

Prof. Dr. Mostafa Elhosseini

Professor of Smart Systems Engineering

<https://youtube.com/drmelhosseini>

| Semi-supervised Learning

- It involves using a combination of a small amount of labeled data and a large amount of unlabeled data to train a model.
- This approach is particularly useful when obtaining labeled data is expensive or labor-intensive, but there is a wealth of unlabeled data available.
- The labeled data provides initial guidance or a starting point for the learning process, while the unlabeled data is used to further refine the model's understanding and performance

| Semi-supervised Learning

- Most semi-supervised learning algorithms are combinations of unsupervised and supervised algorithms. For example, a clustering algorithm may be used to group similar instances together, and then every unlabeled instance can be labeled with the most common label in its cluster. Once the whole dataset is labeled, it is possible to use any supervised learning algorithm.

| Semi-supervised Learning - Example

- Imagine we have a dataset of animal images. Our task is to classify these images into different categories, such as dogs, cats, birds, etc. However, we face a common challenge: while we have a large collection of animal images (say, 10,000 images), only a small fraction of them (like 1,000) are labeled. The labeled images have been identified by experts as belonging to specific animal categories, but labeling the rest is resource-intensive and time-consuming.

| Semi-supervised Learning - Example

- **Initial Training on Labeled Data:** We first use the 1,000 labeled images to train a basic image classification model. This is a supervised learning step where the model learns to identify features and characteristics of different animals from these labeled examples.
- **Applying the Model to Unlabeled Data:** Next, we use the trained model to make predictions on the remaining 9,000 unlabeled images. Even though the model is not fully accurate yet, it can still provide probable labels for these images.

| Semi-supervised Learning - Example

- **Refining the Model:** Now, we use a semi-supervised learning technique. For instance, we might use **self-training**, where the model itself selects the unlabeled images for which it is most confident about its predictions. These high-confidence predictions are then treated as **pseudo-labeled data**.
- **Iterative Learning:** The model is retrained on a combination of the original 1,000 labeled images and the **new pseudo-labeled images**. During this process, the model can refine and improve its learning based on a larger dataset.
- **Validation and Testing:** Finally, the improved model is validated and tested on a separate set of labeled images to evaluate its performance.

| Sentiment Analysis for Online Product Reviews

- Imagine an e-commerce company wants to automatically categorize customer reviews of their products into positive, negative, or neutral sentiments. This kind of analysis is valuable for understanding customer satisfaction and improving products or services. However, while the company has a vast collection of customer reviews, only a small portion of them have been manually categorized by sentiment.

| Sentiment Analysis for Online Product Reviews

- **Initial Training on Labeled Data:** The company starts by training a sentiment analysis model on the **small subset of reviews** that have been **manually labeled**. This initial step is **supervised learning**, where the model learns to identify words, phrases, and patterns associated with positive, negative, and neutral sentiments.
- **Applying the Model to Unlabeled Data:** Next, the trained model is used to predict sentiments for the large set of unlabeled reviews. Although the model's predictions are not fully reliable yet, it can still provide probable sentiment classifications for these reviews.

| Sentiment Analysis for Online Product Reviews

- **Refining the Model:** Here, a semi-supervised learning technique is employed. For instance, the company might use a method like **co-training**, where **two different models are trained separately** on different sets of features and then used to label the unlabeled data for each other.
- **Iterative Learning:** The sentiment analysis model is iteratively retrained on the original labeled dataset plus the newly labeled data from the model's own predictions. This iterative process allows the model to improve its understanding and accuracy.
- **Validation and Testing:** The refined model is then validated and tested on a separate set of manually labeled reviews to assess its accuracy and ability to generalize.

| Customer Segmentation in Retail

- Imagine a **retail company** wants to segment its customers to better understand their shopping habits and preferences. The goal is to use this segmentation to tailor marketing strategies and improve customer service. The company has a large database of customer information, but only a small subset of this data has been manually categorized into segments based on previous marketing campaigns.

| Customer Segmentation in Retail

- Initial Clustering with Unlabeled Data: The company starts by applying a clustering algorithm (like K-means) to the entire dataset of customer information. This step is unsupervised, as it doesn't rely on the previously categorized segments. The algorithm groups customers into clusters based on similarities in their shopping behaviors, demographics, and preferences.
- Labeled Data to Identify Cluster Characteristics: The small subset of customers that have been previously categorized is then used to label the corresponding clusters. For instance, if customers from a particular segment are mostly found in one cluster, that cluster can be labeled accordingly (e.g., "high-value customers", "frequent buyers", etc.).

| Customer Segmentation in Retail

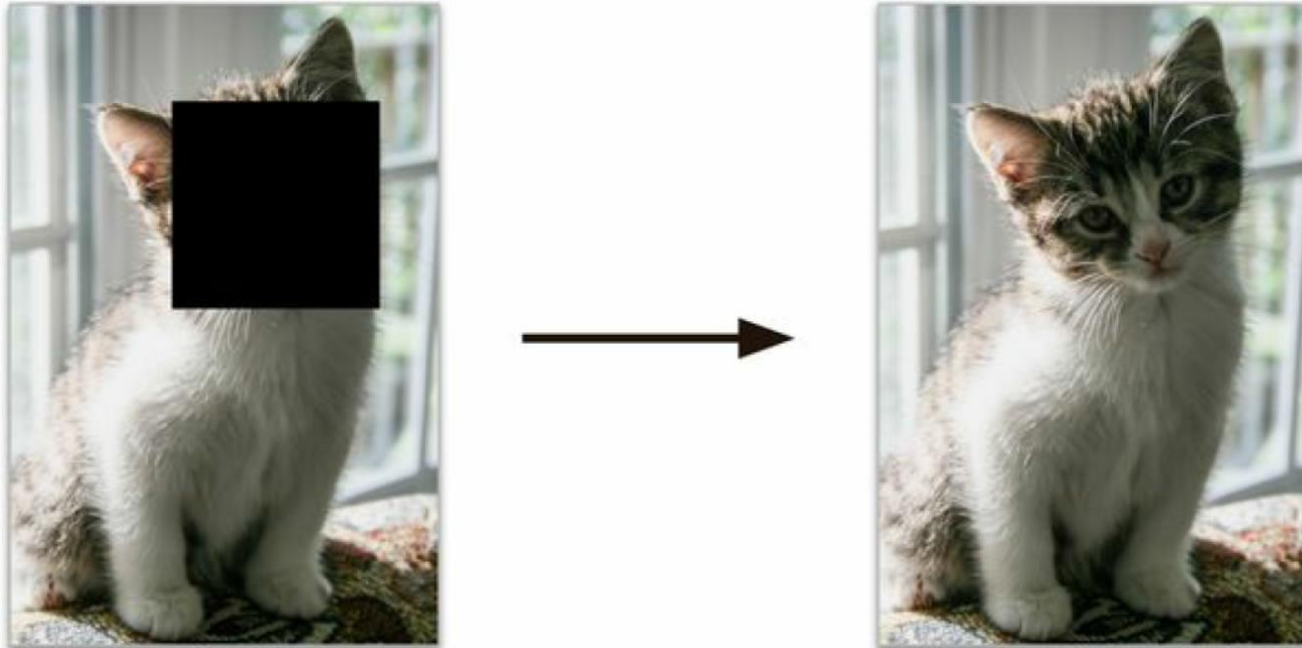
- Refining Clusters with Labeled Data: The labeled clusters are now analyzed to understand their defining characteristics. This information is used to refine the clustering process, possibly by adjusting the number of clusters or the clustering algorithm's parameters.
- Semi-Supervised Learning: The refined clustering model is applied again to the entire dataset. This time, the model has more guidance due to the insights gained from the labeled data. It's an iterative process, where the model alternates between clustering (unsupervised) and adjusting its approach based on the labeled data (semi-supervised).

| Customer Segmentation in Retail

- Validation and Application: Finally, the company validates the resulting customer segments by checking how well they align with known customer behaviors and preferences. The refined segments are then used for targeted marketing campaigns, personalized recommendations, and other customer engagement strategies.

| Self-supervised Learning

- Self-supervised learning is an approach in machine learning, particularly in the field of artificial intelligence (AI), where the system learns to understand and represent data by predicting part of its input from other parts of its input.
- It's a type of unsupervised learning technique that doesn't rely on explicit external labeling of the data. Instead, it generates its own supervisory signals from the input data.



| Self-supervised Learning

- To predict one part of an image given another part, or for text, to predict the next word in a sentence.
- The idea is that by solving these tasks, the model learns useful representations of the data, which can then be used for other tasks such as classification, detection, or segmentation

| Example

- Imagine a project where the goal is to train a model that can understand and analyze images, which could later be used for tasks like object recognition or image classification.
- However, we do not have a dataset of labeled images (images tagged with descriptions or categories). Instead, we have a large collection of unlabeled images.

| Example

- Creating a **Pretext Task**: We set up a pretext task for the model, such as predicting a portion of an image based on the rest of it. For example, we might divide an image into a grid and remove some sections, then ask the model to predict the missing sections based on the remaining parts of the image.
 - In self-supervised learning, the **initial task created by the model** (such as predicting a missing part of the data) is called the **pretext task**. The knowledge gained from this task is then applied to **downstream tasks**, which are the actual tasks of interest (like object recognition in images).

| Example

- **Training the Model on the Pretext Task:** The model is trained on this task using the large collection of unlabeled images. By trying to predict the missing parts of the images, the model learns about important features such as edges, shapes, colors, textures, and how these features are typically arranged in natural images.
- **Learning Feature Representations:** Through this training, the model develops an internal representation of image features. It learns to recognize patterns and structures that are common in the visual world, effectively understanding the 'language' of images.

| Example

- Applying the Learned Representations: The learned feature representations can then be used for downstream tasks like object detection or image classification. Even though the model was not explicitly trained on labeled images for these tasks, the representations it learned are robust and can be fine-tuned with a smaller labeled dataset for specific applications.