

Advancements in Self-Supervised Learning for Visual Representation Learning

^[1] Sachin Ramdas Thorat, ^[2] Dr. Davendranath G Jha

^[1] Department of Data Science and Technology, K.J. Somaiya Institute of Management, Somaiya Vidhyavihar University, Mumbai, India

^[2] Professor Data Science and Technology, K.J. Somaiya Institute of Management Somaiya Vidhyavihar University, Mumbai, India

Email: ^[1] sachin.thorat@somaiya.edu, ^[2] dgjha@somaiya.edu

Orcid id: ^[1] <https://orcid.org/0009-0003-8262-0659>

Abstract— Self-Supervised Learning (SSL) leverages inherent data structures to generate self-supervised learning signals, enabling efficient representation learning from vast unlabeled datasets. This bypasses the need for extensive labeling, unlike supervised learning, and unlocks the potential of massive unlabeled resources.

In recent years, there has been a rapid advancement in SSL methods for visual representation learning. These methods have achieved state-of-the-art results on various tasks, including image classification, object detection, and segmentation.

The advancement in SSL has led to many benefits. First, it has made learning from much larger datasets possible, which can lead to better performance. Second, it has reduced the need for human labelling, which can be expensive and time-consuming. Third, it has made it possible to learn from data that is not easily labeled, such as medical images or natural language.

Advancement in SSL are still ongoing, and many challenges need to be addressed. One challenge is developing SSL methods to learn from more complex data. Another challenge is to develop SSL methods that can be used for more challenging tasks.

Despite these challenges, the advancement in SSL is a promising direction for future research in machine learning. SSL has the potential to revolutionize how machine learning models are trained, and it can be used to solve a wide variety of problems. It also has the potential to make machine learning more accessible and affordable.

This paper surveys the advancements in self-supervised learning for visual representation learning. Through an in-depth analysis of recent research, the paper contrasts various approaches, evaluates their strengths and limitations, and offers insights into their contributions to enhancing visual representation learning. By examining key approaches, data augmentation strategies, and loss functions, the paper assesses the strengths and limitations of each methodology.

This paper also presents a comprehensive literature review, discusses various methodologies, highlights critical observations, and concludes with insights into the future directions of self-supervised learning for visual representation.

Index Terms— Self-Supervised Learning, Visual Representation Learning

I. INTRODUCTION

Visual representation learning is the process of learning low-dimensional representations of visual data that capture the underlying structure of the data. These representations can be used for various tasks, such as image classification, object detection, and segmentation.

Visual representation learning is pivotal in enabling computers to comprehend images effectively. It serves as the foundation for numerous computer vision tasks, enabling machines to understand and interpret the rich information present in images.

Traditionally, visual representation learning has been done using supervised learning. However, supervised learning requires labeled data, which can be expensive and time-consuming to collect.

SSL addresses this challenge by learning representations from unlabeled data.

SSL is a more scalable approach to visual representation learning, as it does not require labeled data.

Self-supervised learning offers a solution by leveraging inherent patterns within the data to create effective representations.

In SSL, the model is trained to predict the input itself or some transformation of it. SSL works by creating a pretext task that can be solved without labels. The pretext task encourages the model to learn meaningful data representations. For example, in the pretext task of predicting the relative position of image patches, the model is trained to predict the order in which two patches were randomly shuffled.

The field of self-supervised learning has emerged as a powerful paradigm, allowing models to learn meaningful representations from unlabeled data. The essence of self-supervised learning lies in its ability to exploit data's inherent structure by designing tasks that require the model to make sense of the data itself. Unlike traditional supervised learning, which relies heavily on labeled examples, self-supervised learning leverages the abundance of unlabeled data available. This shift in the learning paradigm has garnered significant attention and resulted in ground-breaking achievements in

various domains, particularly in the context of visual representation.

This paper embarks on a journey through the recent advancements in self-supervised learning methodologies tailored explicitly for visual representation learning. Also, it delves into recent advancements in self-supervised learning techniques for visual representation, focusing on their methodologies and applications.

II. LITERATURE REVIEW

Self-supervised learning (SSL) has been a rapidly growing area of research in recent years. SSL can learn representations from unlabeled data, which is often much more abundant than labeled data.

One of the earliest works on SSL is the work of (Mishra S, 2020), who proposed the pretext task of predicting the relative position of image patches. This task has been used to learn representations for various downstream tasks, such as image classification and object detection.

The pretext task for SSL was predicting the relative position of image patches. This task was introduced by (Gidaris S & N, 2018) and effective learned representations for image classification.

Another prevalent pretext task is predicting the rotation of an image. This task was introduced (Radford, 2018) as an effective learned representation for object detection.

In recent years, self-supervised learning has emerged as a powerful paradigm for training deep neural networks without requiring extensive labeled data. This approach leverages the inherent structure within the data itself to create supervisory signals, leading to significant advancements in various domains, including computer vision.

The field of self-supervised learning has witnessed significant growth in recent years. One notable approach is **Contrastive Predictive Coding (CPC)** introduced by (Oord A.V, 2018). CPC trains a network to predict future states of a sequence, forcing the model to learn meaningful representations.

More recently, there has been a trend towards using contrastive learning for SSL. Contrastive learning is a technique where a model is trained to distinguish between pairs of images that are similar and pairs of images that are different.

The most popular contrastive learning method is SimCLR, which was introduced by (Chen T., 2020). SimCLR is a simple yet effective method that has been shown to achieve state-of-the-art results on various visual representation learning tasks. SimCLR effectively learns learning representations for various tasks, including image classification, object detection, and segmentation.

SwAV (Huang J, 2021) is another breakthrough, focusing on clustering and grouping similar images in the latent space. **BYOL** (Grill J. B., 2020) improves on this by using two neural networks to learn from each other iteratively. These methodologies highlight the evolution from simple pretext tasks to more complex and effective learning strategies.

Self-supervised learning has been particularly successful in learning meaningful representations from visual data, paving the way for improved performance in downstream tasks such as object recognition, semantic segmentation, and image retrieval. This literature review delves into the latest research on self-supervised learning methods for visual representation learning, highlighting key contributions and insights.

Contrastive Learning Approaches:

- **SimCLR: A Simple Framework for Contrastive Learning of Visual Representations** (Chen T. K., 2020) Introduced the SimCLR framework, which demonstrated the effectiveness of contrastive learning for visual representation. SimCLR maximizes the agreement between differently augmented views of the same image and has shown remarkable performance in various benchmarks.
- **Contrastive Predictive Coding (CPC):** (O, 2020) introduced CPC, which trains a network to predict future representations of a sequence, encouraging the model to capture meaningful features
- **BYOL: Bootstrap Your Own Latent** (Grill J. B., 2020) proposed a novel approach to self-supervised learning using two online neural networks to learn representations. The target network predicts the output of the online network, creating a positive feedback loop that improves representation quality. BYOL employs two neural networks to learn representations by interacting as a target and an online network, eliminating the need for negative samples. BYOL uses online networks to predict the output of target networks, iteratively refining representations. This approach eliminates the need for negative samples in the contrastive loss framework, leading to stable training dynamics.
- **MoCo v3: Revisiting Momentum for Contrastive Learning** (Van Gansbeke W., 2021) extended the Momentum Contrast (MoCo) framework by introducing improvements that enhance its convergence properties and overall performance. MoCo v3 builds upon the success of MoCo and further demonstrates the potential of contrastive learning in visual representation tasks. It constructs a dynamic dictionary of negative samples, enhancing the model's ability to capture fine-grained visual details. (Van Gansbeke W., 2021) Extended the MoCo framework in **MoCo v3**, revisiting the momentum for contrastive learning. MoCo v3 incorporates enhancements such as dynamic sampling of negatives and a memory-bank mechanism, improving performance and training efficiency.

Cluster-Based Approaches:

- **SwAV: Unsupervised Learning of Visual Features by Contrasting Cluster Assignments** (Caron M. T. H., 2021) Proposed SwAV, an approach that contrasts cluster assignments to learn visual features. Unlike pairwise instance comparisons, SwAV focuses on

grouping similar images, showing promise in learning semantically meaningful representations.

SwAV leverages cluster assignments to learn visual representations, achieving state-of-the-art performance without requiring instance-level labels. SwAV is an innovative approach that contrasts cluster assignments rather than pairwise instance comparisons. This method learns features by assigning semantically related images to the same cluster. SwAV demonstrates impressive results by effectively utilizing clustering structures.

Hybrid Approaches:

- **DINO: Emerging Properties in Self-Supervised Vision Transformers** (Caron M. T. H., 2021) introduced DINO, a hybrid approach that combines the strengths of self-supervised learning and vision transformers. DINO uses a teacher-student framework to achieve state-of-the-art performance on various vision tasks.
- **ViT: An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale** (Dosovitskiy A., 2020) Proposed the Vision Transformer (ViT) architecture, which transforms images into sequences of patches and applies transformer-based self-attention mechanisms. ViT showcases the potential of self-attention mechanisms in capturing long-range dependencies in images.

Transfer Learning and Domain Adaptation:

- **UDA: Unsupervised Domain Adaptation for Semantic Segmentation with Noisy Labels** (Zou Y., 2018) Explored the application of self-supervised learning for domain adaptation in semantic segmentation tasks. UDA leverages self-supervision to learn domain-invariant features, improving the generalization of models across different domains.

Listed below is a review of recent research papers that have significantly contributed to the advancements in self-supervised learning for visual representation:

- "Unsupervised Representation Learning by Predicting Image Rotations" by (Gidaris S & N, 2018) This paper introduces a self-supervised task of predicting image rotations, which enforces a neural network to learn meaningful features by capturing object semantics and spatial relationships.
- (He K., 2020) Building on the concept of contrastive learning, this paper presents Momentum Contrast (MoCo), a framework that utilizes a queue to maintain a dynamically changing dictionary of negative examples, leading to improved performance in representation learning.
- (Chen T., 2020) SimCLR proposes a simple yet effective framework for contrastive learning. By maximizing agreement between differently augmented views of the same image, SimCLR achieves state-of-the-art results in self-supervised learning.
- (Grill J. B., 2020) BYOL introduces a method that leverages two online networks to learn representations

without negative samples. This approach demonstrates strong performance and efficiency in representation learning.

- (Caron M. T. H., 2021) This study investigates self-supervised training of Vision Transformers (ViTs) and identifies the emergence of properties like translation invariance, object permanence, and part localization

In recent years, self-supervised learning has witnessed significant progress, leading to a various techniques exploiting different proxy tasks to train models. Notable approaches include:

- **Contrastive Learning:** Methods like (Chen T., 2020) use contrastive loss to bring similar samples closer and push dissimilar samples apart in a learned embedding space. Contrastive learning has been a dominant paradigm in self-supervised visual representation learning. (Chen T., 2020) Introduced a framework that leverages contrastive loss to enhance image representations. It maximizes agreement between differently augmented views of the same image, demonstrating its effectiveness in capturing intricate visual pattern.
- **Temporal Coherence:** Utilizing temporal relationships in videos, methods like T4 and TSN learn representations by predicting temporal order or context.
- **Rotation Prediction:** Models like SwAV (Caron M. M. I., 2020) and BYOL (Grill J. B., 2020) predict image rotations to induce visual understanding and invariances.
- **Context Prediction:** Context-based methods like CPC (Oord A.V, 2018) and CMC (Yunjie Tian, 2021) predict the context of a patch within an image, promoting the capture of high-level semantics.

There has been a lot of recent progress in SSL for visual representation learning. Some of the most notable advances include:

- The development of new pretext tasks that are more effective at learning meaningful representations.
- The use of deep learning architectures that are better suited for SSL.
- The development of new optimization methods that are more efficient for training SSL models.

III. METHODOLOGY

This study, investigates the recent advances in SSL for visual representation learning. We evaluate various SSL methods on multiple tasks and compare their performance to supervised learning methods.

This study investigates the advancements in self-supervised learning for visual representation learning. We evaluated a variety of SSL models on several visual representation learning tasks.

The standard components of self-supervised learning methodologies:

- **Data Augmentation:** Most methods heavily rely on data augmentation techniques, such as random cropping,

flipping, and color jittering, to create diverse views of the same image.

- **Contrastive Loss:** Contrastive loss functions encourage the model to maximize similarity between positive pairs (augmented views of the same image) and minimize similarity with negative pairs (views of different images).
- **Architecture Choices:** Recent trends involve using large-scale neural architectures, often pre-trained on a massive dataset, followed by fine-tuning on a specific task.

Self-supervised learning methods have significantly evolved, driven by innovative strategies and novel architectures. This section delves into the methodologies employed in recent advancements of self-supervised learning for visual representation, encompassing data augmentation techniques, contrastive loss functions, and novel approaches that leverage contextual cues for representation learning.

Data Augmentation Strategies: Data augmentation is a cornerstone of self-supervised learning, creating diverse training samples from a single unlabeled image. Recent research emphasizes the importance of carefully designed augmentations to enable models to capture different aspects of visual content.

- **SimCLR** introduced a set of augmentations, including random cropping, color distortion, and Gaussian blur, to create positive pairs of augmented views. These augmentations encourage the model to learn features invariant to variations in view and enhance its ability to distinguish meaningful patterns in the data (Chen T., 2020).
- **BYOL** employs a "target" and "online" network, each with distinct augmentations. The online network generates positive samples for the target network, promoting alignment of the learned representations (Grill J. B., 2020).

Contrastive Loss Functions: Contrastive loss functions form the crux of self-supervised learning methods. These loss functions encourage the model to differentiate between positive pairs (similar instances) and negative pairs (dissimilar instances). Popular contrastive loss functions include:

- **NT-Xent Loss:** The Normalized Temperature-Scaled Cross-Entropy (NT-Xent) loss (Chen T., 2020) is a widely used contrastive loss that maximizes agreement between positive pairs while penalizing similarity between negative pairs.

- **InfoNCE Loss:** The Information Noise-Contrastive Estimation (InfoNCE) loss is another variant of contrastive loss used in self-supervised learning tasks. It estimates the mutual information between positive pairs.

Contextual Cues and Beyond: Beyond traditional contrastive methods, recent advancements explore new ways of leveraging contextual cues to enhance self-supervised learning.

- **SwAV** proposes a unique approach that contrasts cluster assignments instead of individual instances. This enables the model to learn from the relationships between clusters, leading to semantically meaningful representations (Caron M. T. H., 2021).
- **TransTrack** extends self-supervised learning to video understanding. It introduces temporal cues by tracking object identities across video frames, encouraging the model to learn motion and appearance cues simultaneously (Liang H, 2022)

Contrastive learning approaches like SimCLR and MoCo offer robustness and scalability, while BYOL's simplicity and memory efficiency are noteworthy. Cluster-based methods like SwAV add a new dimension by incorporating semantic relationships. Each methodology has unique strengths and limitations, emphasizing the need for consideration based on the specific context and available resources.

Methodological Comparison: The field of self-supervised learning boasts several methodologies, each with its unique characteristics. A comparison of the latest methodologies reveals the following:

- **Robustness and Generalization:** Methods like SimCLR, with diverse augmentations and large batch sizes, tend to produce robust representations that generalize well to downstream tasks.
- **Semantic Information Capture:** SwAV's emphasis on cluster assignments demonstrates promise in capturing semantically meaningful features by leveraging the inherent structure within data.
- **Ease of Implementation:** BYOL's elimination of negative samples simplifies implementation, potentially reducing the complexity of hyperparameter tuning.

Methodological Advancements: A Comparative Analysis: This section presents a comparison of the latest methodological advancements in self-supervised learning for visual representation, highlighting their strengths and limitations.

Methodology	Approach	Strengths	Limitations	Research Papers
Contrastive Learning Approaches				
SimCLR	Contrastive Learning	Diverse augmentations enhance robustness.	Large batch sizes might require substantial computational resources.	(Chen T., 2020)
		High-quality representations through contrastive loss.	Pair-wise comparisons may not capture higher-order relationships.	
MoCo	Momentum-Based	Momentum-based approach reduces batch size dependence.	Hyper parameter sensitivity may impact performance.	(He K., 2020)

		Queue-based negative samples enhance training efficiency.	Limited data might restrict effectiveness.	
--	--	---	--	--

Augmentation Diversity Strategies				
BYOL	Projection-Based	Eliminates the need for negative samples.	Hyper parameter tuning is critical for success.	(Grill J. B., 2020)
		Consistently improves visual representations.	Convergence may require extensive resources.	
Cluster-Based Approaches				
SwAV	Cluster-Based	Cluster assignments capture semantic relationships.	Performance sensitive to the number of clusters and hyper parameters. Implementation complexity compared to simpler contrastive methods.	(Caron M. M. I., 2020)

The comparison reveals that each methodology has its strengths and limitations. Contrastive methods like SimCLR and MoCo offer robustness, while BYOL and SwAV introduce innovative approaches. The choice of methodology should consider available resources, problem context, and desired outcomes.

The comparison of methodologies in self-supervised learning for visual representation underscores the diversity of approaches available. While each methodology possesses distinct strengths and limitations, they collectively contribute to the advancement of visual representation learning. By carefully selecting and adapting these methodologies, researchers can continue to drive progress in the field

Data Augmentation Strategies: Different data augmentation strategies present distinct trade-offs

Sr#	Augmentation strategies	Strengths	Limitations
1	Random Cropping and Flipping	These strategies introduce variability and encourage invariance to spatial transformations simple yet effective in enhancing robustness	Overly aggressive cropping can remove critical context from the image, affecting the model's understanding
2	Color Distortions	Color distortions enable the model to be less sensitive to lighting variations, contributing to better generalization	Extreme color distortions might lead to unrealistic images that deviate from the natural data distribution
3	Rotation	Predicting rotation angles encourages the model to learn viewpoint-invariant features,	Rotation-based methods might struggle with complex scenes or objects with

		which is essential for various tasks	ambiguous orientations
--	--	--------------------------------------	------------------------

Contrastive Loss Functions: Contrastive loss functions exhibit distinct characteristics:

Sr#	Name	Strengths	Limitations
1	NT-Xent Loss	Encourages clear separation between positive and negative pairs.	Proper selection of the temperature parameter is crucial for optimal performance.
2	InfoNCE Loss	InfoNCE loss estimates mutual information between instances, offering a different perspective on contrastive learning	Estimating mutual information can be computationally intensive, impacting training efficiency.

Evaluation of SSL Models

Evaluation with labels

Self-supervised pretraining is mainly evaluated on image classification since it has been at the core of computer vision for decades. The three main common protocols are k-nearest neighbours (KNN), linear, and full fine-tuning evaluations.

KNN is widely recognized machine learning algorithm and has been extensively used throughout the fields. Regarding image classification, a KNN classifier determines a data point's label from its neighbour's labels. K-NN classifiers are easy to use because they require few adjustments and are fast to implement.

Linear- Linear probing, where a linear classifier is trained on top of pre-trained features, is a popular method for evaluating SSL performance. Feature representations, a.k.a. linear-probing evaluation, was introduced by (Zhang, 2016)

Most of the time, it is done simply by appending a linear layer at the end of the frozen backbone and optimizing its parameters for a few epochs (around 100). Sometimes, as introduced by (Bao H, 2021).

MLP Whereas simple linear probing can be used to analyse SSL representation, a more nuanced approach employing a two- or three-layer multilayer perceptron (MLP) can offer deeper insights into the information captured. (Bordes F, 2023) Present some results that compare different evaluation regimes using a linear or a nonlinear probe. In Figure below, one can observe that it's possible to gain accuracy when using a multilayer layer perceptron instead of a linear probe.

Full Fine-tuning: The Masked Auto-encoders (MAE) paper (Feichtenhofer C, 2022) re-introduced

Fine-tuning is the main evaluation metric. The main arguments are that linear probing is Independent of both fine-tuning and transfer learning.

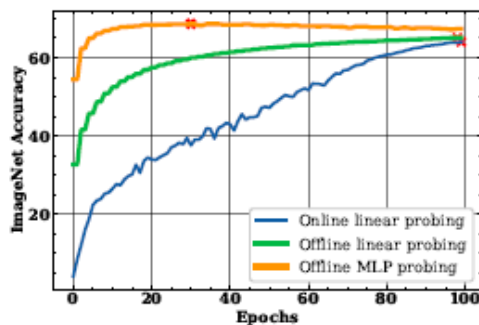


Figure []: Figure from (Bordes F, 2023) Depiction of the classifier probe trained to predict the Imagenet-1k labels from the output of a Resnet50 backbone during SimCLR training (**online**) and post-training (**offline**) using a linear or MLP classifiers.

Evaluation without labels

Using a pretext task such as rotation prediction can facilitate performance evaluation without labels, as demonstrated in (Reed J, 2021) for data augmentation policy selection. This approach requires training the classifier beforehand for the pretext-task and assuming that rotations were not part of the pre-training augmentations or the model would be invariant.

Dataset	Method	VICReg		SimCLR	DINO	
		cov.	inv.	temp.	t-temp.	s-temp.
ImageNet	ImageNet Oracle	68.2	68.2	68.5	72.3	72.4
	α -ReQ	67.9	67.5	63.5	71.7	66.2
	RankMe	67.8	67.9	67.1	72.2	72.4
OOD	ImageNet Oracle	68.7	68.7	68.7	71.9	72.5
	α -ReQ	68.1	67.8	65.1	71.8	68.5
	RankMe	67.7	68.3	67.6	71.8	72.5

Table 3: Hyperparameter selection using the common supervised linear probe strategy

(ImageNet oracle), RankMe and α -ReQ. OOD indicates the average performance over iNaturalist18, Places 205, Sun397, EuroSat, StanfordCars, CIFAR-10, CIFAR-100, Pascal VOC2007, CLEVR-cnt and FOOD101. Without any label, optimization, or parameters, RankMe recovers most of the performance obtained using the ImageNet validation set, highlighting its strength as a hyper parameter selection tool.

Another simple way to evaluate SSL methods, RankMe, was introduced by (Garrido Q, 2022). The idea is to use the effective rank of representations, defined as the entropy of the singular value distribution of the embeddings. It can be computed as:

$$\text{RankMe}(Z) = \exp \left(- \sum_{k=1}^{\min(N,K)} p_k \log p_k \right), \quad p_k = \frac{\sigma_k(Z)}{\|\sigma(Z)\|_1} + \epsilon$$

IV. CONCLUSION

SSL is a promising approach to learning representations from unlabelled data. The advancements in SSL have led to significant improvements in the performance of SSL models on various tasks. We believe that SSL will continue to be a popular research area in the future.

The advancements in self-supervised learning for visual representation have revolutionized the field of computer vision. From simple pretext tasks to complex contrastive learning and clustering strategies, self-supervised methods have demonstrated their ability to extract meaningful features from unlabelled data. The success of these approaches underscores their potential to reduce the need for labelled data, thereby making visual representation learning more accessible and efficient.

Significant progress has been made in the rapidly evolving field of self-supervised learning for visual representation. The various approaches discussed in this paper have demonstrated the potential of leveraging unlabeled data to learn meaningful representations. While these methods have shown promise, challenges such as scalability and fine-tuning on specific tasks remain. As technology advances, self-supervised learning will likely continue to play a pivotal role in pushing the boundaries of visual representation learning.

Methodological advancements in self-supervised learning have significantly impacted the field of visual representation learning. The combination of innovative data augmentation strategies, contrastive loss functions, and network architectures has enabled self-supervised models to excel in various downstream tasks. As methodologies continue to evolve, the future of self-supervised learning holds the promise of even more powerful and versatile visual representations.

With the power to train on vast unlabelled data comes many benefits. While traditional supervised learning methods are trained on a specific task, SSL learns generic representations useful across many tasks.

SSL can be beneficial in domains such as medicine where labels are costly or the specific task cannot be known as a priority (Krishnan R, 2022). There's also evidence that SSL models can learn representations that are more robust to adversarial examples, label corruption, and input perturbations—and are fairer—than to their supervised counterparts (Hendrycks D, 2019).

V. FUTURE DIRECTIONS

Self-supervised learning has emerged as a powerful approach to learning meaningful representations from unlabeled data, alleviating the need for large annotated datasets.

The potential for further advancement in self-supervised learning for visual representation is substantial. Exploring hybrid approaches that combine self-supervision with weak supervision or meta-learning could yield even more impressive results. Additionally, investigating the transferability of learned representations across diverse domains remains an intriguing avenue for research.

Self-supervised learning has opened up new avenues for research and applications in visual representation learning. While the field has made remarkable progress, challenges remain, such as improving the efficiency of training and understanding the generalization capabilities of self-supervised models. Future research should focus on refining these techniques and exploring their potential in real-world scenarios.

The future directions of research in this area include:

- Developing new pretext tasks that are more effective at learning representations.
- Developing new approaches to contrastive learning.

Investigating the use of SSL for other tasks, such as natural language processing and speech recognition.

REFERENCES

- [1] Alexander Kolesnikov, X. Z. (2019). Revisiting Self-Supervised Visual Representation Learning. *Computer Vision and Pattern Recognition*.
- [2] Bao H, D. L. (2021). BEiT: BERT pre-training of image transformers. *arXiv preprint*, 2106.08254.
- [3] Bordes F, B. R. (2023). Towards democratizing joint-embedding self-supervised. *arXiv preprint*, 01986.
- [4] Caron M., M. I. (2020). Unsupervised learning of visual features by contrasting cluster assignment. *Advances in neural information processing systems*, 9912-24.
- [5] Caron M., T. H. (2021). Emerging Properties in Self-Supervised Vision Transformers. In *Proceedings of the IEEE/CVF international conference on computer vision*, 9650-60.
- [6] Chen T., K. S. (2020). Simclr: A simple framework for contrastive learning of visual representations. *ICML'20: Proceedings of the 37th International Conference on Machine Learning*, (pp. 1597-607).
- [7] Chen, T. I. (2020). SimSiam: Self-supervised learning with contrastive learning of siamese networks. *arxiv*.
- [8] Chen, T. K. (2020). A Simple Framework for Contrastive Learning of Visual Representations. *International Conference on Machine Learning*, (pp. 1597-1607). PMLR.
- [9] Chen, T. K. (2022). Revisiting momentum for contrastive learning. *arXiv*.
- [10] Dosovitskiy A., B. L. (2020). Transformers for image recognition at scale. *arXiv*, 11929.
- [11] Feichtenhofer C, F. H. (2022). Masked autoencoders as spatiotemporal learners. *Advances in neural information processing systems*, 35, 35946-58.
- [12] Garrido Q, B. R. (2022). Rankme: Assessing the downstream performance of pretrained self-supervised representations by their rank. In *International Conference on Machine Learning*. 10929-10974.
- [13] Gidaris S, S. P., & N, K. (2018). Unsupervised Representation Learning by Predicting Image Rotations. *International Conference on Learning Representations*, (p. 1803.07728). Paris.
- [14] Grill J. B., S. F. (2020). Bootstrap your own latent-a new approach to self-supervised learning. *Advances in neural information processing systems*, 33, 21271-84.
- [15] He K., F. H. (2020). Momentum contrast for unsupervised visual representation learning. *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, (pp. 9729-9738). 9729-9738.
- [16] He, K. Z. (2016). Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, (pp. 770-778).
- [17] Hendrycks D, M. M. (2019). Using self-supervised learning can improve model robustness and uncertainty. *Advances in neural information processing systems*, 32.
- [18] Huang J, G. S. (2021). Deep Clustering by Semantic Contrastive Learning. *British Machine Vision Conference*.
- [19] Jiangliu Wang, J. J.-H. (2020). Self-supervised Video Representation Learning by Pace Prediction. *ECCV*.
- [20] Krishnan R, R. P. (2022). Self-supervised learning in medicine and healthcare. *Nature Biomedical Engineering*, 6(12), 1346-52.
- [21] Maregu Assefa, W. J. (2022). Self-Supervised Multi-Label Transformation Prediction for Video Representation Learning. *Journal of Circuits, Systems and Computers*.
- [22] Mishra S, A. S. (2020). Learning visual representations for transfer learning by suppressing texture. *ArXiv*, 01901.
- [23] O, H. (2020). Data-efficient image recognition with contrastive predictive coding. In *International conference on machine learning*. PMLR.
- [24] Oord A.V, L. Y. (2018). Representation learning with contrastive predictive coding. *ArXiv*, 1807.03748.
- [25] Prathamesh Sonawane, S. D. (2021). Self-Supervised Visual Representation Learning Using Lightweight Architectures. *Arxiv*.
- [26] Priya Goyal, D. M. (2019). Scaling and Benchmarking Self-Supervised Visual Representation Learning. *IEEE International Conference on Computer Vision*.
- [27] Radford, A. N. (2018). Improving language understanding by generative pre-training. *arXiv preprint*.
- [28] Reed J, M. S. (2021). Self Augment: Automatic augmentation policies for self-supervised learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, (pp. 2674 -2683).
- [29] Sravanti Addepalli, K. B. (2022). Towards Efficient and Effective Self-Supervised Learning of Visual Representations. *European Conference on Computer Vision*.
- [30] Tengda Han, W. X. (2020). Self-supervised Co-training for Video Representation Learning. *Neural Information Processing Systems*.
- [31] Terbouche, H. a. (2022). Comparing Learning Methodologies for Self-Supervised Audio-Visual Representation Learning. *IEEE Access*, 41622-41638.
- [32] Van Gansbeke W., V. S. (2021). Revisiting contrastive methods for unsupervised learning of visual representations. *Advances in Neural Information Processing Systems*, 34, 16238-50.
- [33] Xiao Zhang, M. M. (2020). Self-Supervised Visual Representation Learning from Hierarchical Grouping. *Neural Information Processing Systems*.

- [34] Yunjie Tian, L. X. (2021). Semantic-Aware Generation for Self-Supervised Visual Representation Learning. arXiv.org.
- [35] Zhang, X. B. (2016). Improved evaluation of unsupervised image representations using linear probing. Advances in neural information processing systems, 3766 -3744.
- [36] Zou Y., Y. Z. (2018). Unsupervised domain adaptation for semantic segmentation via class-balanced self-training. In Proceedings of the European conference on computer vision, 289-305.
- [37] Thorat, S. R., Jha, D. G., Sharma, A. K., & Katkar, D. V. (2024). Wrist fracture detection using self-supervised learning methodology. Journal of Musculoskeletal Surgery and Research, 8(2), 133-141.



IFERP[®]
Explore Your Research Journey...