

# Momentum Contrast for Unsupervised Visual Representation Learning

Sowe Ebou A<sup>1\*</sup> and Bah Yunusa A<sup>2</sup>

<sup>1</sup>*School of Computer Science and Artificial Intelligence, Wuhan University of Technology, China*

<sup>2</sup>*Department of Geography, Ohio University, China*

## \*Corresponding Author

Sowe Ebou A, School of Computer Science and Artificial Intelligence, Wuhan University of Technology, China.

**Submitted:** 2025, Jan 27; **Accepted:** 2025, Mar 07; **Published:** 2025, Mar 14

**Citation:** Ebou, S. A., Yunusa, B. A. (2025). Momentum Contrast for Unsupervised Visual Representation Learning. *Eng OA*, 3(3), 01-06.

## Abstract

*This brief report presents a novel unsupervised learning representation learning method called momentum contrast. Momentum contrast uses a contrastive learning technique to learn representations by comparing features of related yet dissimilar images for efficient feature extraction and unsupervised representation learning. Similar images are grouped together, and dissimilar images are placed far apart. The method builds upon previous works in contrastive learning but includes a momentum optimisation step to improve representation learning performance and generate better quality representations. Experiments on various datasets demonstrate that momentum contrast is able to learn high-quality representations, allowing us to directly use them to achieve competitive performance with fewer labelled examples.*

**Keywords:** Unsupervised Learning, Contrastive Learning, Visual Representation

## 1. Introduction

Visual representation learning is a crucial component of many computer vision applications. In recent years, there has been a growing interest in unsupervised methods for learning visual representations. Unsupervised methods do not require labelled data, making them more versatile and applicable to a wider range of tasks. One unsupervised method that has gained popularity in recent years is Momentum Contrast (MoCo) for unsupervised visual representation learning. MoCo is a mechanism for building dynamic dictionaries for contrastive learning and can be used with various pretext tasks [1]. In this report, I will explore the concept of momentum contrast for unsupervised visual representation learning and its performance in comparison to other unsupervised learning methods. I will also explore some experimental results of MoCo, based on these experiments, I will drive some conclusions.

## 2. Background Study

Traditionally, supervised learning has been the dominant paradigm for training deep neural networks for visual recognition tasks. Contrastive learning (CL) is one of the prominent keystones of self-supervised learning. It fosters discriminability in the representation [2-6]. However, there are several limitations to this approach. Firstly, obtaining large labelled datasets can be expensive and time-consuming, especially for specialized tasks. Secondly, even with large labelled datasets, the resulting models may not generalize well to new, unseen images. Finally, supervised

learning is not applicable to many domains where labelled data is scarce or unavailable. Unsupervised visual representation learning involves training a neural network to learn a set of features that can be used to represent images in a meaningful way. These features can then be used for a variety of tasks, such as image classification, object detection, and semantic segmentation. Unsupervised learning methods typically rely on data augmentation techniques to generate a large amount of diverse training data. The goal is to train a neural network to learn a set of features that are invariant to these augmentations. The core idea is to pull representations of “similar” images (referred to as positives) close while “dissimilar” images (negatives) are contrasted in feature space. Such methods implemented this idea using an instance discrimination pretext task where only transformed versions of the same images are taken as positives while augmented versions of other images are negatives [7].

## 3. Methodology

### 3.1. Contrastive Learning as Dictionary Look-up

Contrastive learning is a machine learning technique that aims to learn useful representations by contrasting pairs of examples. Contrastive learning can drive a variety of pretext tasks [1].

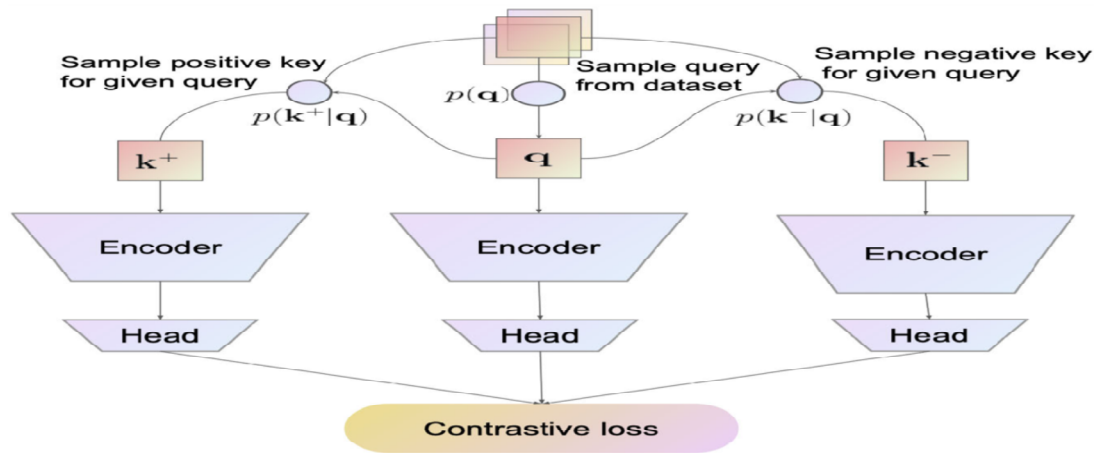
Even though contrastive learning has become prominent in recent years due to the success of large pre-trained models in the fields of natural language processing (NLP) and computer vision (CV),

the seminal idea dates back at least to the 1990s [8,9]. MoCo uses contrastive learning technique by making the dynamic dictionary large and consistence. Among the most successful of the recent self-supervised approaches to learning visual representations, a subset of these termed “contrastive” learning methods have achieved the most success [10].

The negative samples used for contrastive learning are obtained from a dynamic dictionary. Initially, the dictionary is empty. As the model learns, the encoder representations of the input data are stored in the dictionary. The dictionary is maintained using a

queue-based mechanism, where new representations replace the oldest ones. This way, the dictionary captures a wide range of negative samples over time, providing diverse negative pairs for contrastive learning.

By continuously updating the dictionary of negative samples and training the encoder using the contrastive loss, MoCo encourages the model to capture useful and semantically meaningful representations. The dynamic nature of the dictionary allows the model to adapt to changing data distributions and learn robust representations.



**Figure 1:** Overview of the Contrastive Representation Learning framework. Its components are: a similarity and dissimilarity distribution to sample positive and negative keys for a query, one or more encoders and transform heads for each data modality and a contrastive loss function evaluate a batch of positive and negative pairs [10].

### 3.2. Momentum Contrast

Momentum Contrast (MoCo) is an unsupervised learning method that was introduced in a paper by He et al. in 2019. The method is based on the idea of using a momentum encoder to generate a set of target features. During training, the momentum encoder is updated using a moving average of the weights of the online encoder. The online encoder is trained to generate features that are similar to the target features.

The MoCo method has several advantages over other unsupervised learning methods. One advantage is that it is computationally efficient, allowing for larger batch sizes and longer training times. Another advantage is that it is more effective at learning representations that are invariant to data augmentations. This is achieved by using a larger set of augmentations during training. MoCo can outperform its supervised pre-training counterpart in 7 detection/segmentation tasks on PASCAL VOC, COCO, and other datasets, some- times surpassing it by large margins [1].

MoCo v1 has attracted significant attention by demonstrating superior performance over supervised pre-training counterparts in downstream tasks while making use of large negative samples, decoupling the need for batch size by introducing a dynamic dictionary [1,11].

### 3.3. Approach

The Momentum Contrast (MoCo) approach is a recent development in unsupervised representation learning that has shown state-of-the-art performance on a variety of visual recognition tasks. The MoCo approach is based on the principle of contrastive learning, which has been shown to be effective for unsupervised representation learning. The basic idea of contrastive learning is to learn representations that are invariant to certain transformations (e.g., rotations, translations, etc.) while maintaining discriminative power for similar images.

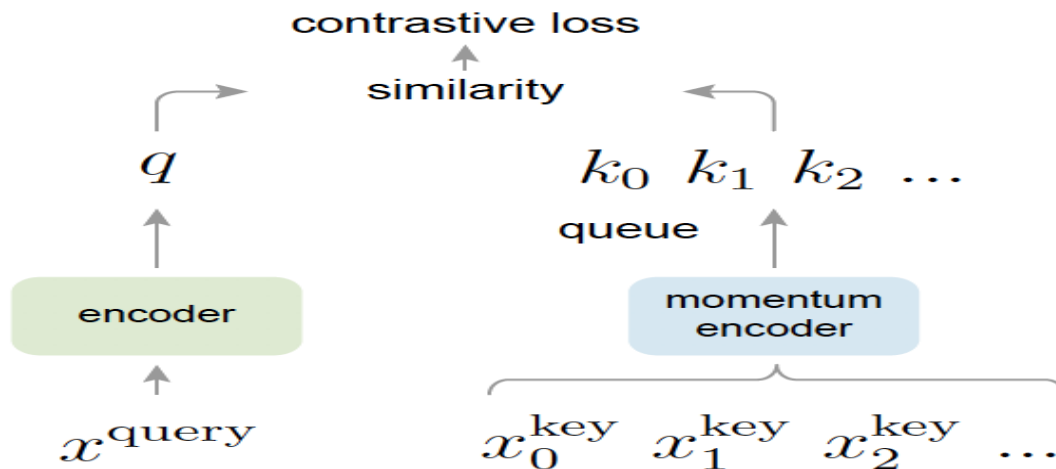
The MoCo approach builds on the contrastive learning principle by introducing a momentum-based update rule that improves the stability and convergence of the training process. Specifically, the MoCo approach uses a memory bank to store a large number of negative examples that are used to compute contrastive losses during training. The memory bank is updated using a momentum-based update rule that averages the parameters of the current model with those of a slowly-updated "queue" model. This update rule helps to stabilize the training process by providing a more consistent source of negative examples.

### 3.4. Architecture

The MoCo architecture consists of two main components: an encoder network and a memory bank. The encoder network is a

deep convolution neural network (CNN) that is trained to extract features from raw image data. The memory bank is a large matrix that stores a set of negative examples that are used to compute contrastive losses during training. The memory bank is updated using a momentum-based update rule that averages the parameters of the current model with those of a slowly-updated "queue" model.

The encoder network consists of a series of convolutional layers followed by a global average pooling layer and a fully connected layer. The output of the fully connected layer is a vector of fixed length that represents the image features. The encoder network is trained using a contrastive loss function that encourages similar images to have similar feature representations, while dissimilar images have different feature representations.



**Figure 2:** Momentum Contrast (MoCo) trains a visual representation encoder by matching an encoded query  $q$  to a dictionary of encoded keys using a contrastive loss [1]. The dictionary keys  $\{k_0, k_1, k_2, \dots\}$  are defined on-the-fly by a set of data samples [1].

### 3.4.1. Dictionary as a Queue

Momentum contrast uses a dictionary queue encoded with keys. The queue is updated by adding the representation of the current image to the queue and removing the oldest representation. The dictionary acts as a "memory bank" that stores a history of feature representations. The queue is updated using a momentum-based update rule, which allows the model to maintain a smooth and stable representation of the feature space.

During training, the images are split into two groups: a query group and a key group. The query group is used to compute a query feature representation, while the key group is used to compute a set of key feature representations. The query feature representation is then compared to the key feature representations stored in the dictionary using a contrastive loss function.

### 3.4.2. Momentum Update

Momentum update is a key component in Momentum Contrastive learning, a technique commonly used in self-supervised learning tasks such as image or video representation learning. It helps improve the stability and convergence speed of the learning process by introducing a momentum term during the update of the model's parameters.

In MoCo, the momentum update is used to update the model's parameters based on the current gradient and the momentum term. The momentum update can be visualized as follows:

Initialize the model's parameters and momentum parameters.

**At each training iteration:**

a). Compute the gradients of the loss function with respect to the

parameters using the current mini-batch of data.

b). Update the momentum parameters using the momentum update equation:  $\mathbf{v}_t = \alpha * \mathbf{v}_{t-1} + (1 - \alpha) * \mathbf{g}_t$ , where  $\mathbf{v}_t$  is the velocity term at time step  $t$ ,  $\alpha$  is the momentum coefficient, and  $\mathbf{g}_t$  is the gradient at time step  $t$ .

c). Update the model's parameters using the momentum parameters:  $\theta_{t+1} = \theta_t + \mathbf{v}_t$ , where  $\theta_{t+1}$  represents the updated parameters at time step  $t+1$ .

The momentum update equation calculates the velocity term by combining the previous velocity  $\mathbf{v}_{t-1}$  and the current gradient  $\mathbf{g}_t$ . The momentum coefficient  $\alpha$  determines the contribution of the previous velocity compared to the current gradient. A higher  $\alpha$  value gives more weight to the previous velocity, resulting in a smoother and more stable update trajectory.

The momentum update allows the model to accumulate information from previous gradients and helps the optimization process by maintaining a consistent direction of updates. This can help the model escape shallow local minima and converge faster to better representations.

Keep in mind that while the momentum update is an essential component of MoCo, the specific implementation details and hyperparameters may vary depending on the exact architecture and training setup.

## 4. Performance

MoCo has been shown to outperform other unsupervised learning methods on several benchmark datasets, including ImageNet,

---

CIFAR-10, and CIFAR-100. MoCo achieves state-of-the-art performance on these datasets without using any labelled data. MoCo has also been shown to be effective at learning representations for downstream tasks such as object detection and semantic segmentation.

#### 4.1. Experiment Results

Momentum Contrast (MoCo) has shown impressive results in various experimental settings and benchmark datasets. MoCo's success can be attributed to its innovative contrastive learning framework, which encourages the model to learn discriminative representations by contrasting positive and negative samples. Positive samples in MoCo are augmented versions of the same image, while negative samples are drawn from a queue. This learning approach enables the model to pull similar samples closer together in the learned representation space while pushing dissimilar samples apart.

Here are some notable experimental results achieved by MoCo:

##### 4.1.1. ImageNet Classification

MoCo achieved state-of-the-art performance on the ImageNet-1K dataset, which consists of 1.28 million labelled images spanning 1,000 object categories. In the MoCo v2 paper, the authors reported top-1 accuracy of 60.6% using ResNet-50, surpassing previous self-supervised methods and approaching the performance of supervised methods.

In the MoCo v2 paper, the authors reported impressive results using the MoCo framework with the ResNet-50 architecture. They achieved a top-1 accuracy of 60.6% on the ImageNet-1K dataset. This performance surpassed previous self-supervised methods and approached the performance of supervised methods, which rely on human-labelled data.

This achievement is significant because it demonstrates the effectiveness of self-supervised learning approaches like MoCo in learning high-quality representations from large amounts of unlabelled data. By leveraging the power of contrastive learning and momentum encoders, MoCo was able to capture meaningful visual features that improved image classification accuracy.

The ability of MoCo to achieve competitive results on the challenging ImageNet-1K dataset indicates that self-supervised learning has the potential to bridge the gap between supervised and unsupervised methods. It opens up possibilities for utilizing vast amounts of unannotated data to learn representations that approach the performance of supervised models, reducing the reliance on human-labelled data.

These advancements in self-supervised learning and the success of MoCo on ImageNet-1K have contributed to the growing interest and exploration of self-supervised methods in various computer vision tasks and domains.

##### 4.1.2 Transfer Learning

MoCo has demonstrated strong transfer learning capabilities.

Pretrained models using MoCo representations have been successfully transferred to various downstream tasks such as object detection, semantic segmentation, and instance segmentation. By initializing the models with MoCo pretrained weights, significant performance gains have been observed compared to training from scratch.

##### 4.1.3. Few-Shot Learning

MoCo has also shown promise in few-shot learning scenarios, where the goal is to recognize novel classes with limited labelled examples. By leveraging the learned representations, MoCo has been used as a feature extractor to achieve competitive performance on few-shot learning benchmarks like miniImageNet and tieredImageNet.

##### 4.1.4. Robustness to Adversarial Attacks

MoCo has demonstrated improved robustness to adversarial attacks compared to supervised learning. By training on large-scale unlabelled data, MoCo learns more generalizable representations that are less susceptible to adversarial perturbations.

Adversarial attacks involve making intentional and often imperceptible modifications to input data in order to deceive a machine learning model. These perturbations can lead to incorrect predictions or misclassification. Adversarial attacks are a significant concern in various domains, including computer vision.

By training on large amounts of unlabelled data, MoCo learns to capture underlying patterns and structures in the data that are more resilient to adversarial perturbations. The robustness stems from the model's ability to generalize across a diverse set of samples and learn more invariant representations. This generalizability helps in reducing the vulnerability to adversarial attacks.

Furthermore, the contrastive learning framework of MoCo, where positive samples are augmented versions of the same image and negative samples are drawn from a queue, encourages the model to pull similar samples closer together while pushing dissimilar samples apart. This contrastive objective promotes the learning of discriminative features that are less susceptible to adversarial perturbations.

##### 4.1.5. Generalization

MoCo has been shown to generalize well across different domains and datasets. For example, models pretrained on ImageNet using MoCo have been successfully transferred to domain-specific datasets such as Pascal VOC and COCO, achieving competitive performance.

These results highlight the effectiveness of MoCo in learning meaningful and transferable image representations without relying on explicit labels, thereby enabling broader applications and reducing the need for large amounts of labelled data.

#### 5. Conclusion

Momentum Contrast is an effective unsupervised learning method for visual representation learning. It has several advantages over

---

other unsupervised learning methods, including computational efficiency and better invariance to data augmentations. MoCo has been shown to achieve state-of-the-art performance on several benchmark datasets, making it a promising approach for unsupervised visual representation learning. Based on this, I drive the following conclusions:

### 5.1. Improved Self-supervised Learning

MoCo has shown significant improvements over traditional self-supervised learning methods. By leveraging a momentum encoder, MoCo creates a dynamic and consistent queue of negative samples, enabling better learning of representations without the need for manual annotations. It addresses the limitations of traditional self-supervised learning approaches by introducing a momentum encoder and a dynamic queue of negative samples.

Overall, MoCo's use of a momentum encoder and a dynamic queue of negative samples has led to significant improvements in self-supervised learning. By leveraging these techniques, MoCo has demonstrated state-of-the-art performance on various benchmark datasets, surpassing previous methods that relied on manual annotations or supervised learning.

### 5.2. Used Contrastive Learning

MoCo utilizes a contrastive learning framework, where positive samples are augmented versions of the same image and negative samples are drawn from the queue. This encourages the model to pull positive samples together while pushing away negative samples, leading to more discriminative representations.

The main idea behind MoCo is to encourage the model to pull positive samples (augmented versions of the same image) closer together in the embedding space while pushing negative samples (drawn from a queue of other images) further apart. This helps in learning more discriminative and meaningful representations.

The contrastive learning process in MoCo can be summarized in the following steps:

**Online Encoder and Target Encoder:** MoCo maintains two encoders, the online encoder and the target encoder. The target encoder represents a slowly moving average of the online encoder's weights, which provides a consistent and stable set of representations.

**Positive Pair Augmentation:** To create positive pairs, an image is randomly augmented multiple times. These augmented versions of the same image serve as positive samples. By applying different augmentations, the model learns to capture different views of the same underlying object or scene.

**Negative Sample Selection:** Negative samples are drawn from a queue that stores representations of other images in the dataset. The queue acts as a source of negative samples that the model should be pushed away from. This helps in learning more robust and discriminative representations.

**Contrastive Loss:** MoCo uses a contrastive loss function to train the model. The contrastive loss encourages the model to maximize the similarity between positive pairs while minimizing the similarity between positive and negative pairs. This loss formulation drives the model to learn representations that are more discriminative and generalize well to downstream tasks.

By training with the contrastive learning framework of MoCo, the model can learn powerful representations from unlabeled data, which can then be fine-tuned or transferred to supervised tasks with limited labeled data, leading to improved performance.

### 5.3. Transferability

MoCo has demonstrated excellent transferability of learned representations. Pretrained models using MoCo have been shown to achieve state-of-the-art performance on various downstream tasks such as image classification, object detection, and semantic segmentation. This indicates that the learned representations capture general visual concepts that can be transferred across different tasks.

The pretrained models from MoCo provide a strong starting point for fine-tuning on specific supervised tasks with limited labelled data. By leveraging the knowledge acquired during the unsupervised pretraining phase, these models can effectively generalize and adapt to new tasks. This transfer learning approach saves significant computational resources and reduces the need for extensive labelled data.

The success of MoCo and similar self-supervised learning methods highlights the potential of unsupervised learning in capturing meaningful representations that benefit a wide range of downstream applications.

### 5.4. Robustness to Label Noise

MoCo's self-supervised nature makes it robust to label noise in the training data. Since it does not rely on human annotations, MoCo can learn from large amounts of unlabelled data, which is often easier to obtain compared to accurately labelled data. This is particularly advantageous in scenarios where labelled data is scarce or expensive.

Label noise refers to errors or inconsistencies in the annotations of the training data. In traditional supervised learning, these errors can negatively impact the model's performance as it learns from mislabelled examples. However, self-supervised learning methods like MoCo bypass the need for explicit labels by utilizing pretext tasks that create supervision signals from the data itself.

By leveraging unlabelled data, MoCo can learn powerful representations that are robust to label noise. The model learns to capture inherent patterns and structure in the data, allowing it to generalize well even in the presence of noisy or imperfect labels. This is particularly valuable in real-world scenarios where obtaining accurately labelled data can be challenging, such as in large-scale datasets or domains where expert annotations are

scarce.

Furthermore, the ability of MoCo to learn from unlabelled data makes it highly scalable. It can leverage vast amounts of readily available unlabelled data, such as images or text corpora, enabling the model to learn rich and meaningful representations without the need for manual annotations. This scalability makes MoCo an attractive approach in situations where obtaining labelled data is limited or costly.

Overall, MoCo's self-supervised learning paradigm empowers the model to learn robust representations from unlabelled data, making it particularly advantageous in scenarios with label noise, scarce labelled data, or when access to accurately labelled data is challenging.

### 5.5. Promising Applications

The effectiveness of MoCo in visual representation learning opens up opportunities for various applications. It can be applied to domains such as computer vision, robotics, and autonomous systems, where understanding visual information is crucial for perception, decision-making, and action.

In conclusion, Momentum Contrast (MoCo) has emerged as a powerful technique for visual representation learning. Its ability to learn from large-scale unlabeled data, robustness to label noise, and excellent transferability make it a valuable tool for advancing computer vision research and applications [12].

### References

1. He, K., Fan, H., Wu, Y., Xie, S., & Girshick, R. (2020). Momentum contrast for unsupervised visual representation learning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition* (pp. 9729-9738).
2. Federici, M., Dutta, A., Forré, P., Kushman, N., & Akata, Z. (2020). Learning robust representations via multi-view information bottleneck. *arXiv preprint arXiv:2002.07017*.
3. Misra, I., Zitnick, C. L., & Hebert, M. (2016). Shuffle and learn: unsupervised learning using temporal order verification. In *Computer Vision–ECCV 2016: 14th European Conference, Amsterdam, The Netherlands, October 11–14, 2016, Proceedings, Part I 14* (pp. 527-544). Springer International Publishing.
4. Schroff, F., Kalenichenko, D., & Philbin, J. (2015). Facenet: A unified embedding for face recognition and clustering. In *Proceedings of the IEEE conference on computer vision and pattern recognition* (pp. 815-823).
5. Sohn, K. (2016). Improved deep metric learning with multi-class n-pair loss objective. *Advances in neural information processing systems*, 29.
6. Wang, X., & Gupta, A. (2015). Unsupervised learning of visual representations using videos. In *Proceedings of the IEEE international conference on computer vision* (pp. 2794-2802).
7. Chen, X., Fan, H., Girshick, R., & He, K. (2020). Improved baselines with momentum contrastive learning. *arXiv preprint arXiv:2003.04297*.
8. Zhang, Y., Hu, X., Sapkota, N., Shi, Y., & Chen, D. Z. (2022, December). Unsupervised feature clustering improves contrastive representation learning for medical image segmentation. In *2022 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)* (pp. 1820-1823). IEEE.
9. Becker, S., & Hinton, G. E. (1992). Self-organizing neural network that discovers surfaces in random-dot stereograms. *Nature*, 355(6356), 161-163.
10. Bromley, J., Guyon, I., LeCun, Y., Säckinger, E., & Shah, R. (1993). Signature verification using a "siamese" time delay neural network. *Advances in neural information processing systems*, 6.
11. Le-Khac, P. H., Healy, G., & Smeaton, A. F. (2020). Contrastive representation learning: A framework and review. *Ieee Access*, 8, 193907-193934.
12. Nguyen, T., Pham, T. X., Zhang, C., Luu, T. M., Vu, T., & Yoo, C. D. (2023). DimCL: Dimensional contrastive learning for improving self-supervised learning. *IEEE Access*, 11, 21534-21545.

**Copyright:** ©2025 Sowe Ebou A, et al. This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.