

Pruning the Paradox: How CLIP’s Most Informative Heads Enhance Performance While Amplifying Bias

Avinash Madasu[♡]
[♡]Intel Labs

avinash.madasu@intel.com

Vasudev Lal^{♣*}
[♣]Oracle

Phillip Howard^{◇*}
[◇]Thoughtworks

phillip.howard@thoughtworks.com

Abstract

CLIP is one of the most popular foundation models and is heavily used for many vision-language tasks, yet little is known about its inner workings. As CLIP is increasingly deployed in real-world applications, it is becoming even more critical to understand its limitations and embedded social biases to mitigate potentially harmful downstream consequences. However, the question of what internal mechanisms drive both the impressive capabilities as well as problematic shortcomings of CLIP has largely remained unanswered. To bridge this gap, we study the conceptual consistency of text descriptions for attention heads in CLIP-like models. Specifically, we propose Concept Consistency Score (CCS), a novel interpretability metric that measures how consistently individual attention heads in CLIP models align with specific concepts. Our soft-pruning experiments reveal that high CCS heads are critical for preserving model performance, as pruning them leads to a significantly larger performance drop than pruning random or low CCS heads. Notably, we find that high CCS heads capture essential concepts and play a key role in out-of-domain detection, concept-specific reasoning, and video-language understanding. Moreover, we prove that high CCS heads learn spurious correlations which amplify social biases. These results position CCS as a powerful interpretability metric exposing the paradox of performance and social biases in CLIP models.

1 Introduction

Large-scale vision-language (VL) models such as CLIP (Radford et al., 2021) have significantly advanced state-of-the-art performance in vision tasks in recent years. Consequently, CLIP has been extensively used as a foundational model for downstream tasks such as video retrieval, image generation, and segmentation (Luo et al., 2022; Liu

et al., 2024; Brooks et al., 2023; Esser et al., 2024; Kirillov et al., 2023). This has enabled the construction of compositional models combining CLIP with other foundation models, thereby increasing the functionality of CLIP while also adding complexity to the overall model structure. However, as these models gain prominence in real-world applications, their embedded social biases (Howard et al., 2024; Hall et al., 2023; Seth et al., 2023) have emerged as a critical concern with potentially harmful consequences. Despite the growing body of work documenting these biases, a fundamental question remains: *what mechanisms within these models’ architectures drive both their impressive capabilities and problematic shortcomings?*

Recent interpretability advances (Gandelsman et al.) have made initial progress by decomposing CLIP’s image representations into contributions from individual attention heads, identifying text sequences that characterize different heads’ semantic roles. However, this approach provides only a partial view into CLIP’s inner workings, leaving a critical missing piece: systematic understanding of the visual concepts encoded at the attention head level—and how these concepts underpin both the model’s strengths and its social failures.

Our work addresses this critical gap through a novel interpretability framework we call “conceptual consistency”. This framework systematically analyzes which visual concepts are learned by individual attention heads and how consistently these concepts are processed throughout the model’s architecture. First, we identify interpretable structures within the individual heads of the last four layers of the model using a set of text descriptions. To accomplish this, we employ the TEXTSPAN algorithm (Gandelsman et al.), which helps us find the most appropriate text descriptions for each head. After identifying these text descriptions, we assign labels to each head representing the common property shared by the descriptions. This label-

*Work partially completed while at Intel Labs.

ing process is carried out using in-context learning with ChatGPT. We begin by manually labeling five pairs of text descriptions and their corresponding concept labels, which serve as examples. These examples are then used to prompt ChatGPT to assign labels for the remaining heads.

Leveraging the resulting text descriptions and concept labels of attention heads, we introduce the Concept Consistency Score (CCS), a new interpretability metric that quantifies how strongly individual attention heads in CLIP models align with specific concepts. Using GPT-4o, Gemini and Claude as automatic judges, we compute CCS for each head and classify them into high, moderate, and low categories based on defined thresholds. A key contribution of our work is our targeted soft-pruning experiments which show that heads with high CCS are essential for maintaining model performance; pruning these heads causes a significantly larger performance drop compared to pruning any other heads. We also show that high CCS heads are not only crucial for general vision-language tasks but are especially important for out-of-domain detection and targeted concept-specific reasoning. Additionally, our experiments in video retrieval highlight that high CCS heads are equally vital for temporal and cross-modal understanding. Moreover, we demonstrate that high CCS heads often encode spurious correlations, contributing to social biases in CLIP models. Selective pruning of these heads can reduce such biases without the need for fine-tuning. Together, these results expose a fundamental paradox: while high-CCS heads are indispensable for strong model performance, they are simultaneously key contributors to undesirable biases.

2 Related Work

Early research on interpretability primarily concentrated on convolutional neural networks (CNNs) due to their intricate and opaque decision-making processes (Zeiler and Fergus, 2014; Selvaraju et al., 2017; Simonyan et al., 2014; Fong and Vedaldi, 2017; Hendricks et al., 2016). More recently, the interpretability of Vision Transformers (ViT) has garnered significant attention as these models, unlike CNNs, rely on self-attention mechanisms rather than convolutions. Researchers have focused on task-specific analyses in areas such as image classification, captioning, and object detection to understand how ViTs process and interpret visual in-

formation (Dong et al., 2022; Elguendouze et al., 2023; Mannix and Bondell, 2024; Xue et al., 2022; Cornia et al., 2022; Dravid et al., 2023). One of the key metrics used to measure interpretability in ViTs is the attention mechanism itself, which provides insights into how the model distributes focus across different parts of an image when making decisions (Cordonnier et al., 2019; Chefer et al., 2021). This has led to the development of techniques that leverage attention maps to explain ViT predictions. Early work on multimodal interpretability, which involves models that handle both visual and textual inputs, probed tasks such as how different modalities influence model performance (Cao et al., 2020; Madasu and Lal, 2023) and how visual semantics are represented within the model (Hendricks and Nematzadeh, 2021; Lindström et al., 2021). Aflalo et al. (Aflalo et al., 2022) explored interpretability methods for vision-language transformers, examining how these models combine visual and textual information to make joint decisions. Similarly, Stan et al. (Stan et al., 2024) proposed new approaches for interpreting vision-language models, focusing on the interactions between modalities and how these influence model predictions. Our work builds upon and leverages the methods introduced by Gandselman et al., 2024) to interpret attention heads, neurons, and layers in vision-language models, providing deeper insights into their decision-making processes. Several works (Shu et al., 2023; Mayilvahanan et al.; Moeller et al., 2024) explored the out-of-domain generalization of CLIP models on several benchmarks. In contrast to these works, our work provides a new interpretability metric to understand the decision making processes of attention heads in CLIP models.

3 Quantifying interpretability in CLIP models

3.1 Preliminaries

In this section, we describe our methodology, starting with the TEXTSPAN (Gandselman et al.) algorithm and its extension across all attention heads in multiple CLIP models using in-context learning. Specifically, we deal with attention heads in the image encoder. TEXTSPAN associates each attention head with relevant text descriptions by analyzing the variance in projections between head outputs and candidate text representations. Through iterative projections, it identifies distinct components aligned with different semantic aspects. While ef-

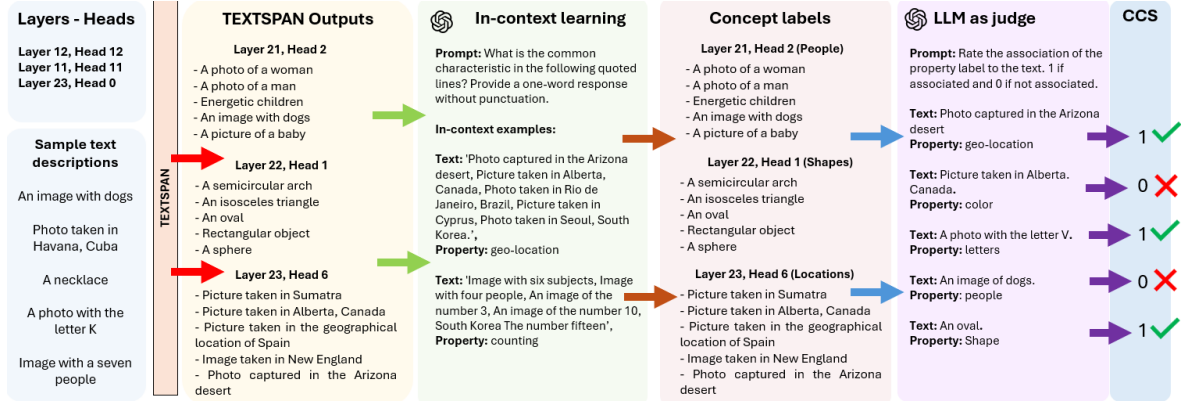


Figure 1: Illustration of our approach to computing Concept Consistency Score for each attention head.

High CCS ($CCS = 5$)	Moderate CCS ($CCS = 3$)	Low CCS ($CCS \leq 1$)
L23.H11 ("People")	L23.H0 ("Material")	L21.H6 ("Professions")
Playful siblings	Intricate wood carving	Photo taken in the Italian pizzerias
A photo of a young person	Nighttime illumination	thrilling motorsport race
Image with three people	Image with woven fabric design	Urban street fashion
A photo of a woman	Image with shattered glass reflections	An image of a Animal Trainer
A photo of a man	A photo of food	A leg
L22.H10 ("Animals")	L11.H0 ("Locations")	L10.H6 ("Body parts")
Image showing prairie grouse	Photo taken in Monument Valley	A leg
Image with a donkey	Majestic animal	colorful procession
Image with a penguin	An image of Andorra	Contemplative monochrome portrait
Image with leopard print patterns	An image of Fiji	Graceful wings in motion
detailed reptile close-up	Image showing prairie grouse	Inviting reading nook
L23.H5 ("Nature")	L11.H11 ("Letters")	L9.H2 ("Textures")
Intertwined tree branches	A photo with the letter J	Photo of a furry animal
Flowing water bodies	A photo with the letter K	Closeup of textured synthetic fabric
A meadow	A swirling eddy	Eclectic street scenes
A smoky plume	A photo with the letter C	Serene beach sunset
Blossoming springtime blooms	awe-inspiring sky	Minimalist white backdrop

Table 1: Examples of high, moderate and low CCS heads.

Model	High	Medium	Low
ViT-B-32-OpenAI	8	13	4
ViT-B-32-datacomp	11	9	8
ViT-B-16-OpenAI	10	14	4
ViT-B-16-LAION	13	12	6
ViT-L-14-OpenAI	16	13	11
ViT-L-14-LAION	21	14	3

Table 2: Count of high, medium and low CCS heads in CLIP models.

effective at linking heads to descriptive text spans, TEXTSPAN does not assign explicit concept labels. We inherited the set of 3498 text spans from Gandelsman et al. (Gandelsman et al.), which were generated by prompting ChatGPT-3.5 to produce general image descriptions using the prompts shown below. After obtaining an initial set, ChatGPT was manually prompted to generate more examples of specific patterns we found (e.g., image with a lunar eclipse, Oceanic coral reef). In the

next section, we detail our method for labeling the concepts learned by individual CLIP heads.

3.2 Concept Consistency Score (CCS)

We introduce the Concept Consistency Score (CCS) as a systematic metric for analyzing the concepts (properties) learned by transformer layers and attention heads in CLIP-like models. This score quantifies the alignment between the textual representations produced by a given head and an assigned concept label. Existing interpretability methods (Chefer et al., 2021; Abnar and Zuidema, 2020; Barkan et al., 2021) are largely descriptive and do not identify which internal components (e.g., attention heads) are reliably responsible for model behavior. This lack of discriminative insight limits their utility in scenarios where intervention or control is required. CCS addresses this by measuring concept consistency across heads, enabling us to isolate specialized heads that consistently encode particular concepts. This makes interpretability

Model	CIFAR-10						CIFAR-100					
	Original	K=5	K=7	K=9	K=11	K=13	Original	K=5	K=7	K=9	K=11	K=13
ViT-B-32-OpenAI	75.68	71.31	71.86	73.66	74.53	70.54	65.08	56.07	62.21	60.17	56.42	59.05
ViT-B-32-datacomp	72.07	70.50	70.59	69.61	69.93	70.20	54.95	53.14	53.59	53.73	53.81	53.91
ViT-B-16-OpenAI	78.10	63.93	71.83	76.08	64.67	76.74	68.22	51.70	52.45	61.90	62.59	55.54
ViT-B-16-LAION	82.82	78.91	81.25	80.25	78.03	78.56	76.92	65.55	69.94	72.46	67.59	69.49
ViT-L-14-OpenAI	86.94	86.29	85.44	85.89	85.49	85.58	78.28	75.66	74.85	74.70	76.73	77.19
ViT-L-14-LAION	88.29	86.48	85.99	86.34	87.28	85.99	83.37	80.07	80.75	82.03	82.14	80.11

Table 3: Accuracy (%) on CIFAR-10 and CIFAR-100 datasets across different values of TEXTSPAN (K).

Model	Kappa	SC (ρ)	Kendall (τ)
ViT-B-32-OpenAI	0.821	0.737	0.781
ViT-B-16-LAION	0.813	0.773	0.737
ViT-L-14-OpenAI	0.827	0.751	0.758

Table 4: Results between human judgment and LLM judgment on CCS labelling. SC denotes Spearman’s correlation.

actionable by supporting tasks like pruning, debugging, or understanding failure modes with certainty and control. Figure 1 illustrates our approach, with the following sections detailing each step in computing CCS.

3.2.1 Extracting Text Representations

From each layer and attention head of the CLIP model, we obtain a set of five textual outputs, denoted as $\{T_1, T_2, T_3, T_4, T_5\}$, referred to as TEXTSPANS. These outputs serve as a textual approximation of the concepts encoded by the head.

3.2.2 Assigning Concept Labels

Using in-context learning with ChatGPT, we analyze the set of five TEXTSPAN outputs and infer a concept label C_h that best represents the dominant concept captured by the attention head h . This ensures that the label is data-driven and reflects the most salient pattern learned by the head.

3.2.3 Evaluating Concept Consistency

To assess the consistency of a head with respect to its assigned concept label, we employ three state-of-the-art foundational models, GPT-4o, Gemini 1.5 pro and Claude Sonnet as external evaluators. For each TEXTSPAN T_i associated with head h , GPT-4o determines whether it aligns with the assigned concept C_h . The Concept Consistency Score (CCS) for head h is then computed as:

$$CCS(h) = \sum_{i=1}^5 \mathbb{I}[T_i \text{ aligns with } C_h]$$

where $\mathbb{I}[\cdot]$ is an indicator function that returns 1 if T_i to be consistent with C_h , and 0 otherwise. To ensure a high standard of reliability, we define consistency strictly—only if all three LLM judges independently rate T_i as consistent with C_h . This requirement for unanimous agreement minimizes the influence of individual model biases or variability in judgment (Liu and Zhang, 2025), thereby enhancing the robustness and trustworthiness of the overall concept consistency score.

We define $CCS@K$ as the fraction of attention heads in a CLIP model that have a Concept Consistency Score (CCS) of K . This metric provides a global measure of how many heads strongly encode interpretable concepts. A higher $CCS@K$ value indicates that a greater proportion of heads exhibit strong alignment with a single semantic property. Mathematically, $CCS@K$ is defined as:

$$CCS@K = \frac{1}{H} \sum_{h=1}^H \mathbb{I}[CCS(h) = K]$$

where H is the total number of attention heads in the model, $CCS(h)$ is the Concept Consistency Score of head h , $\mathbb{I}[\cdot]$ is an indicator function that returns 1 if $CCS(h) = K$, and 0 otherwise. This metric helps assess the overall interpretability of the model by quantifying the proportion of heads that consistently capture well-defined concepts. Table 1 shows the examples of heads with different CCS scores.

Next, we categorize each attention head based on its Concept Consistency Score (CCS) into three levels: high, moderate, and low. A head is considered to have a high CCS if all of its associated text descriptions align with the labeled concept, indicating that the head is highly specialized and likely encodes features relevant to that concept. Moderate CCS heads exhibit partial alignment, with three out of five text descriptions matching the concept label, suggesting that they capture the concept to some extent but not exclusively. In contrast, low

CCS heads have zero or only one matching description, implying minimal relevance and indicating that these heads are largely unrelated to the given concept. This categorization provides insight into the degree of concept selectivity exhibited by individual attention heads. Table 1 shows examples of different types of CCS heads and Table 2 shows the count of high, moderate and low CCS heads in all CLIP models.

3.3 Evaluating LLM Judgment Alignment with Human Annotations

In the previous section, we introduced the Concept Consistency Score (CCS), computed using three LLM judges as an external evaluator. This raises an important question: *Are LLM evaluations reliable and aligned with human assessments?* To investigate this, we conducted a human evaluation study comparing LLM-generated judgments with human annotations. We selected 100 TEXTSPAN descriptions from three different models, along with their assigned concept labels, and asked one of the authors to manually assess the semantic alignment between each span and its corresponding label.

Table 4 reports the agreement metrics between human and LLM evaluations, including Cohen’s Kappa, Spearman’s ρ , and Kendall’s τ . The Kappa values exceed 0.8, indicating extremely substantial agreement, while the correlation scores consistently surpass 0.7, confirming strong alignment. These results validate the use of LLMs as reliable evaluators in concept consistency analysis. The high agreement with human judgments suggests that LLMs can effectively assess semantic coherence, offering a scalable alternative to manual annotation. In the next section, we introduce the tasks and datasets used in our experiments.

3.4 Sensitivity of CCS to the number of TEXTSPAN descriptions

Next, we analyze the robustness of CCS computation to number of TEXTSPAN (K) descriptions. To directly address this issue, we conduct experiments evaluating CCS performance across different numbers of TextSpan descriptions on CIFAR-10 and CIFAR-100 datasets. The results are depicted in the table 3. From the results, it is evident that performance remains remarkably stable across all K values, with variations typically within 1-3% accuracy points across multiple CLIP models. Importantly, the relative performance ordering between different model configurations remains consistent regardless

of K , indicating that our concept alignment assessment captures meaningful patterns. Performance does not consistently improve with higher K values, suggesting that five TEXTSPAN descriptions capture sufficient conceptual diversity without unnecessary computational overhead. These findings directly validate that CCS is robust to the number of TEXTSPAN descriptions and not overly sensitive to this hyperparameter choice.

3.5 Experimental Setting

3.5.1 Tasks

For our experiments, we use a wide variety of datasets focusing on the tasks of image classification, out-of-domain classification, video retrieval, and bias measurement. Below, we mention datasets for each of the tasks.

Image classification: CIFAR-10 (Krizhevsky et al., 2009), CIFAR-100 (Krizhevsky et al., 2009), Food-101 (Bossard et al., 2014), Country-211 (Radford et al., 2021) and Oxford-pets (Parkhi et al., 2012).

Out-of-domain classification: Imagenet-A (Hendrycks et al., 2021b) and Imagenet-R (Hendrycks et al., 2021a).

Video retrieval: MSRVT (Xu et al., 2016), MSVD (Chen and Dolan, 2011), DiDeMo (Anne Hendricks et al., 2017).

Bias: FairFace (Karkkainen and Joo, 2021), SocialCounterFactuals (Howard et al., 2024).

3.5.2 Models

For experiments we use the following six foundational image-text models: ViT-B-32, ViT-B-16 and ViT-L-14 pretrained from OpenAI-400M (Radford et al., 2021) and LAION2B (Schuhmann et al., 2022). Next, we discuss in detail the results from the experiments.

4 Results and Discussion

4.1 Interpretable CLIP Models: The Role of CCS.

In this section we examine the role of the Concept Consistency Score (CCS) in revealing CLIP’s decision-making process, focusing on the question: *How does CCS provide deeper insights into the functional role of individual attention heads in influencing downstream tasks?* To explore this, we perform a soft-pruning analysis by zeroing out attention weights of heads with extreme CCS values—specifically, high CCS (CCS = 5) and low

Model	CIFAR-10			CIFAR-100			FOOD-101		
	Original	High CCS	Low CCS	Original	High CCS	Low CCS	Original	High CCS	Low CCS
ViT-B-32-OpenAI	75.68	71.31	73.61	65.08	56.07	62.39	84.01	73.42	82.12
ViT-B-32-datacomp	72.07	70.50	70.43	54.95	53.14	53.72	41.66	38.13	40.77
ViT-B-16-OpenAI	78.10	63.93	76.44	68.22	51.70	65.38	88.73	76.35	87.36
ViT-B-16-LAION	82.82	78.91	75.38	76.92	65.55	72.51	86.63	67.54	81.4
ViT-L-14-OpenAI	86.94	86.29	85.97	78.28	75.66	77.55	93.07	90.75	92.79
ViT-L-14-LAION	88.29	86.48	88.19	83.37	80.07	83.25	91.02	86.45	90.35

Table 5: Accuracy comparison of various CLIP models on CIFAR-10, CIFAR-100 and FOOD-101 datasets. The values represent original accuracy, performance after pruning high-CCS heads, and performance after pruning low-CCS heads.

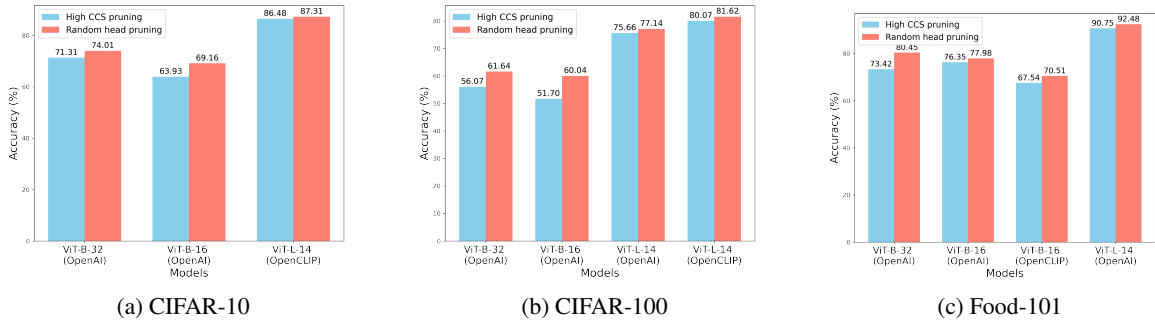


Figure 2: Zero-shot performance comparison for CIFAR-10, CIFAR-100, and Food-101 datasets under different pruning strategies. For random pruning, results are averaged across five runs.

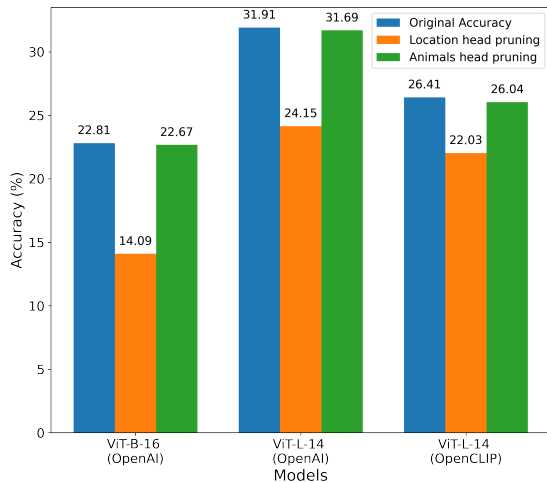


Figure 3: Zero-shot results on Country-211 (location) dataset.

CCS ($CCS \leq 1$). This approach disables selected heads without modifying the model architecture. As shown in Table 5, pruning high-CCS heads consistently causes significant drops in zero-shot classification performance across CIFAR-10, CIFAR-100 and FOOD-101 while pruning low-CCS heads has a minimal effect. This performance gap demonstrates that CCS effectively identifies heads encoding critical, concept-aligned information, making

it a reliable tool for interpreting CLIP’s internal decision-making mechanisms.

We further observe notable variations in pruning sensitivity across model architectures. ViT-B-16 models suffer the most from high-CCS head pruning, implying a reliance on a smaller number of specialized heads. In contrast, ViT-L-14 models show greater resilience, suggesting more distributed representations. Among smaller models, OpenAI-trained models experience larger performance drops than OpenCLIP models when high-CCS heads are pruned. However, in larger models like ViT-L-14, OpenCLIP variants show a slightly higher degradation. These patterns reveal that CCS not only identifies functionally important heads but also captures model-specific and training-specific differences in how conceptual knowledge is organized and utilized within CLIP architectures.

4.2 Pruning equal number of high CCS, low CCS and random heads

In the previous section, we showed that attention heads with high Concept Consistency Scores (CCS) are crucial to CLIP’s performance. To validate whether these heads are truly more important than others, we perform a controlled comparison against random pruning. Specifically, we randomly prune

Model	Country-211			Oxford-pets			ImageNet-A			ImageNet-R		
	Original	High CCS	Low CCS	Original	High CCS	Low CCS	Original	High CCS	Low CCS	Original	High CCS	Low CCS
ViT-B-32-OpenAI	17.16	11.46	16.3	50.07	46.66	48.96	31.49	20.24	28.72	69.09	54.47	64.45
ViT-B-32-datacomp	4.43	4.37	4.37	26.48	25.98	25.33	4.96	4.59	4.65	34.06	31.6	32.47
ViT-B-16-OpenAI	22.81	10.72	21.79	52.72	49.12	51.89	49.85	25.49	47.27	77.37	55.52	74.84
ViT-B-16-LAION	20.45	7.49	16.87	65.79	48.48	49.81	37.97	25.27	27.44	80.56	66.32	71.73
ViT-L-14-OpenAI	31.91	23.21	30.63	61.79	62.04	62.08	70.4	68.15	69.2	87.87	86.56	86.97
ViT-L-14-LAION	26.41	16.38	25.66	54.1	56.12	57.16	53.8	42.44	52.93	87.12	82.22	86.94

Table 6: **Accuracy comparison of various CLIP models on Country-211, Oxford-pets, ImageNet-A and ImageNet-R datasets. The values represent original accuracy, performance after pruning high-CCS heads, and performance after pruning low-CCS heads.**

Model	CIFAR-100			FOOD-101			Country-211		
	Original	High CCS	Low CCS	Original	High CCS	Low CCS	Original	High CCS	Low CCS
ViT-B-32-OpenAI	65.08	62.39	58.58	84.01	82.12	73.24	17.16	16.3	12.23
ViT-B-32-datacomp	54.95	53.72	53.77	41.66	40.77	39.31	4.43	4.37	4.41
ViT-B-16-OpenAI	68.22	65.38	56.17	88.73	87.36	83.23	22.81	21.79	21.74
ViT-B-16-LAION	76.92	72.51	72.73	86.63	81.4	80.51	20.45	16.87	14.39
ViT-L-14-OpenAI	78.28	77.55	73	93.07	92.79	90.15	31.91	30.63	25.38
ViT-L-14-LAION	83.37	83.25	83.35	91.02	90.35	90.25	26.41	25.66	25.43

Table 7: **Accuracy comparison of various CLIP models on CIFAR-100, FOOD-101 and Country-211 datasets. The values represent original accuracy, performance after pruning equal number of high-CCS and low-CCS heads.**

the same number of attention heads—excluding high-CCS heads—and repeat this across five seeds, averaging the results. As illustrated in Figure 2, pruning high-CCS heads consistently causes a significantly larger drop in zero-shot accuracy compared to random pruning across datasets and model variants. In contrast, random pruning results in only minor performance degradation, highlighting the functional importance of high-CCS heads. Interestingly, we also find that larger CLIP models show a smaller performance gap between high-CCS and random pruning, suggesting that larger architectures may be more robust due to greater redundancy or more distributed representations.

Similarly, we conducted experiments where we pruned an equal number of high and low CCS attention heads across multiple datasets (CIFAR-100, FOOD-101, Country-211). Results are shown in the table 7. From the table, we observe that pruning high-CCS heads leads to a substantially larger performance drop, even when the number of pruned heads is held constant. This effectively rules out the explanation that the observed degradation is merely due to pruning more heads in the high-CCS condition. Taken together, these findings support CCS as a reliable and interpretable metric for identifying concept-relevant heads and offer deeper insights into how CLIP organizes conceptual information.

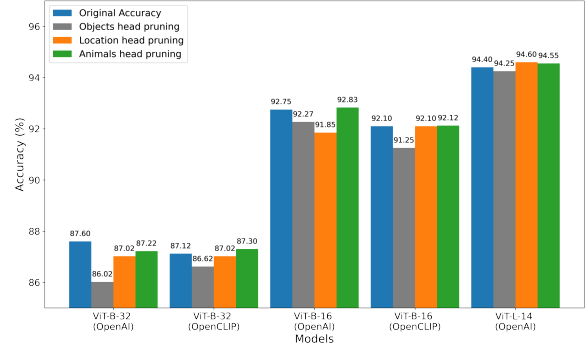


Figure 4: **Zero-shot results on CIFAR-10 (Objects) dataset.**

4.3 High CCS heads are crucial for out-of-domain (OOD) detection

While our earlier experiments primarily focused on in-domain datasets such as CIFAR-10 and CIFAR-100 to validate the Concept Consistency Score (CCS), understanding model behavior under out-of-domain (OOD) conditions is a critical step toward evaluating models’ robustness. Table 6 demonstrates the results on ImageNet-A and ImageNet-R datasets respectively. From the table, we observe that pruning heads with high CCS scores leads to a substantial degradation in model performance, underscoring the critical role these heads play in the model’s decision-making process. Notably, the

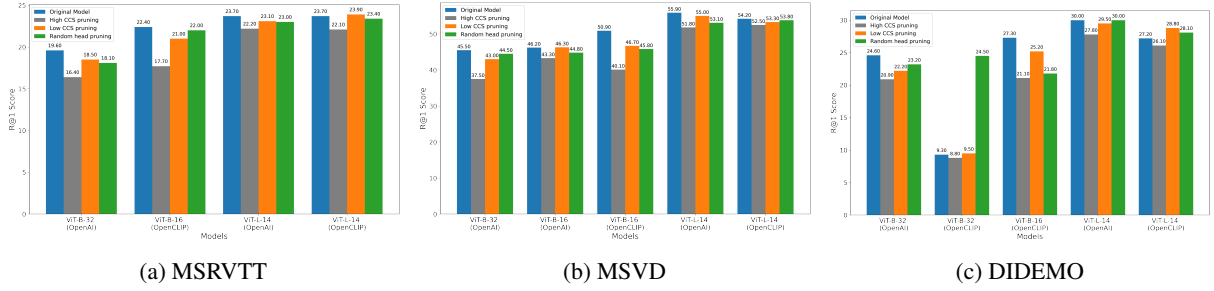


Figure 5: Zero-shot performance comparison of unpruned (original) model, pruning high CSS, low CSS and random heads on video retrieval task.

ViT-B-16-OpenAI model exhibits the most pronounced drop in performance upon pruning high CCS heads, suggesting that this model relies heavily on a smaller set of concept-specific heads for robust feature representation consistent with the observations previously. These results demonstrate that CCS is a powerful metric for identifying attention heads that encode essential, generalizable concepts in CLIP models.

4.4 High CCS heads are crucial for concept-specific tasks.

To investigate the functional role of high Concept Consistency Score (CCS) heads, we conduct concept-specific pruning experiments. In these experiments, we prune heads with high CCS scores corresponding to a target concept (e.g., locations) and evaluate the model’s performance on tasks aligned with that concept, such as location classification. In contrast, we also prune heads associated with unrelated concepts (e.g., animals) and assess the resulting impact on task performance. Our results indicate that pruning high CCS heads leads to a significant drop in task performance, validating that these heads encode essential concept-relevant information. For instance, in the ViT-B-16 model, pruning location heads results in a substantial decrease in location classification accuracy from 22.81% to 14.09%, as shown in Figure 3. Conversely, pruning heads corresponding to unrelated concepts has little effect on performance, demonstrating the concept-specific nature of high CCS heads, as illustrated in Figure 4.

In more general classification tasks, object-related heads consistently exhibit a greater impact on performance than location or color heads. For example, in the ViT-B-32 model, pruning object-related heads leads to a more noticeable accuracy drop (from 87.6% to 86.02%) compared to pruning location or color heads, which result in

smaller reductions (87.02% and 87.22%, respectively). This underscores the greater importance of object-related features in vision tasks. Larger models, such as ViT-L-14, demonstrate a more robust performance to pruning, with smaller accuracy drops when pruning concept-specific heads, suggesting that these models employ more distributed and redundant representations. For instance, pruning object-related heads in ViT-L-14 reduces accuracy only marginally, from 92.1% to 91.25%, with negligible effects from pruning location and color heads. These results not only confirm the effectiveness of CCS as an interpretability tool but also show that high CCS heads are critical for concept-aligned tasks and provide significant insights into how concepts are represented within CLIP-like models.

4.5 Impact of CCS pruning on zero-shot video retrieval.

To further assess the importance of high CCS heads for downstream tasks, we conducted zero-shot video retrieval experiments on three popular datasets (MSRVTT, MSVD, and DIDEMO) under different pruning strategies. The results in Figure 5 show that pruning high CCS (Concept Consistency Score) heads consistently leads to a substantial drop in performance across all datasets, demonstrating their critical role in preserving CLIP’s retrieval capabilities. For instance, on MSRVTT and MSVD, high CCS pruning significantly underperforms compared to low CCS and random head pruning, which show much milder performance degradation. Interestingly, low CCS and random head pruning maintain performance much closer to the original unpruned model, indicating that not all attention heads contribute equally to model competence. This consistent trend across datasets highlights that heads with high CCS scores are essential for encoding concept-aligned information necessary for accurate zero-shot video retrieval.

Model	Race ($\downarrow$)			Gender ($\downarrow$)		
	Original	High CCS	Low CCS	Original	High CCS	Low CCS
ViT-B-32-OpenAI	61.8	60.5	65.0	41.2	41.1	42.5
ViT-B-32-datacomp	49.2	48.3	49.7	21.7	21	21.6
ViT-B-16-OpenAI	64.6	55.6	66.6	40.2	38.2	40.5
ViT-B-16-LAION	59.2	56.7	56.6	43.6	43.1	43.8
ViT-L-14-OpenAI	59.3	59.8	59.9	34.7	32.2	35.2
ViT-L-14-LAION	61.9	55.6	59.7	43	39.2	43.9

Table 8: **Comparison of original and high-CCS pruning on FairFace dataset for race and gender..** We used MaxSkew (K=900) as the metric.

Model	Race ($\downarrow$)			Gender ($\downarrow$)		
	Original	High CCS	Low CCS	Original	High CCS	Low CCS
ViT-B-32-OpenAI	4.1	1.2	2.4	3.7	2.4	1.2
ViT-B-16-OpenAI	0.8	2.0	1.2	2.4	1.2	2.1
ViT-L-14-OpenAI	2.0	0.8	2.0	2.4	2.0	1.6

Table 9: **Comparison of original and high-CCS soft-pruning on SocialCounterFactuals dataset for race and gender..** We used MaxSkew (K=12 for race, K=4 for gender) as the metric.

4.6 CLIP’s high-CCS heads encode features that drive social biases.

Previously, we established that high-CCS heads in CLIP models are crucial for image and video tasks and pruning them leads to significant drop in performance. Now, we investigate if these high CCS heads learn spurious features leading to social biases. For this, we perform soft pruning experiment on FairFace and SocialCounterFactuals datasets. Here given neutral text prompts of 104 occupations*, we measure MaxSkew across race and gender in the datasets. Tables 8 and 9 show the results on FairFace and SocialCounterFactuals datasets respectively.

On the FairFace dataset, pruning high-CCS heads consistently reduces the MaxSkew values for both race and gender across all models. For example, in the ViT-B-16-OpenAI model, pruning high-CCS heads drops the race MaxSkew from 64.6 to 55.6 and the gender MaxSkew from 40.2 to 38.2. Similar reductions are observed across all ViT-B and ViT-L variants. These drops, although modest in some cases, indicate a consistent trend: high-CCS heads are contributing disproportionately to skewed model predictions. The effect is even more evident on the SocialCounterFactuals dataset, where MaxSkew values drop substantially upon pruning high-CCS heads. For instance, in

*List of occupations and prompts can be found in Appendix

ViT-B-32-OpenAI, the race MaxSkew falls from 4.1 to 1.2, and gender MaxSkew from 3.7 to 2.4. Similar reductions occur for other ViT variants, with some pruned models showing more than 50% decrease in bias.

These results reveal a fundamental paradox at the heart of CLIP models: high-CCS heads, while critical for strong performance in tasks such as classification, retrieval, and concept alignment, are also the primary contributors to social bias. Pruning these heads leads to a notable reduction in model bias, as shown in our experiments, but also comes at the cost of reduced performance—a clear trade-off between fairness and utility. This performance-bias paradox underscores the complex role of high-CCS heads: they are both enablers of semantic understanding and carriers of learned stereotypes. The CCS metric, in this context, provides a valuable lens for navigating this tension. It not only aids in interpreting model behavior but also offers a lightweight intervention—soft-pruning—that mitigates bias without requiring expensive fine-tuning.

5 Conclusion

In this work, we proposed Concept Consistency Score (CCS), a novel interpretability metric that quantifies how consistently individual attention heads in CLIP-like models align with semantically meaningful concepts. Through extensive soft-pruning experiments, we demonstrated that heads with high CCS are essential for maintaining model performance, as their removal leads to substantial performance drops compared to pruning random or low CCS heads. Our findings further highlight that high CCS heads are not only critical for standard vision-language tasks but also play a central role in out-of-domain detection and concept-specific reasoning. Moreover, experiments on video retrieval tasks reveal that high CCS heads are crucial for capturing temporal and cross-modal relationships, underscoring their broad utility in multimodal understanding. In addition, we demonstrated that high-CCS heads learn spurious correlations leading to social biases and pruning them mitigates that harmful behaviour without the need for further fine-tuning. Thus, CCS provides an holistic view of interpretability proving the paradox performance vs social biases in CLIP.

6 Limitations

In this work, we experimented primarily on CLIP models. Although CCS metric established the fundamental paradox of performance vs social biases we haven't proved for other vision language models. Hence, we leave extending for more vision language models for future work. Another limitation is the use of LLM models for concept labelling and judging which requires robust manual verification to limit any inconsistencies. Hence, scaling our work to much bigger models with more layers and heads can be a limitation.

References

- Samira Abnar and Willem Zuidema. 2020. Quantifying attention flow in transformers. *arXiv preprint arXiv:2005.00928*.
- Estelle Aflalo, Meng Du, Shao-Yen Tseng, Yongfei Liu, Chenfei Wu, Nan Duan, and Vasudev Lal. 2022. VI-interpret: An interactive visualization tool for interpreting vision-language transformers. In *Proceedings of the IEEE/CVF Conference on computer vision and pattern recognition*, pages 21406–21415.
- Lisa Anne Hendricks, Oliver Wang, Eli Shechtman, Josef Sivic, Trevor Darrell, and Bryan Russell. 2017. Localizing moments in video with natural language. In *Proceedings of the IEEE international conference on computer vision*, pages 5803–5812.
- Oren Barkan, Edan Hauer, Avi Caciularu, Ori Katz, Itzik Malkiel, Omri Armstrong, and Noam Koenigstein. 2021. Grad-sam: Explaining transformers via gradient self-attention maps. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*, pages 2882–2887.
- Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. 2014. Food-101—mining discriminative components with random forests. In *Computer vision—ECCV 2014: 13th European conference, Zurich, Switzerland, September 6–12, 2014, proceedings, part VI 13*, pages 446–461. Springer.
- Tim Brooks, Aleksander Holynski, and Alexei A Efros. 2023. Instructpix2pix: Learning to follow image editing instructions. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 18392–18402.
- Jize Cao, Zhe Gan, Yu Cheng, Licheng Yu, Yen-Chun Chen, and Jingjing Liu. 2020. Behind the scene: Revealing the secrets of pre-trained vision-and-language models. In *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part VI 16*, pages 565–580. Springer.
- Hila Chefer, Shir Gur, and Lior Wolf. 2021. Transformer interpretability beyond attention visualization. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 782–791.
- David Chen and William B Dolan. 2011. Collecting highly parallel data for paraphrase evaluation. In *Proceedings of the 49th annual meeting of the association for computational linguistics: human language technologies*, pages 190–200.
- Jean-Baptiste Cordonnier, Andreas Loukas, and Martin Jaggi. 2019. On the relationship between self-attention and convolutional layers. *arXiv preprint arXiv:1911.03584*.
- Marcella Cornia, Lorenzo Baraldi, and Rita Cucchiara. 2022. Explaining transformer-based image captioning models: An empirical analysis. *AI Communications*, 35(2):111–129.
- Bowen Dong, Pan Zhou, Shuicheng Yan, and Wangmeng Zuo. 2022. Towards class interpretable vision transformer with multi-class-tokens. In *Chinese Conference on Pattern Recognition and Computer Vision (PRCV)*, pages 609–622. Springer.
- Amil Dravid, Yossi Gandelsman, Alexei A. Efros, and Assaf Shocher. 2023. Rosetta neurons: Mining the common units in a model zoo. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 1934–1943.
- Sofiane Elguendouze, Adel Hafiane, Marcilio CP de Souto, and Anaïs Halftermeyer. 2023. Explainability in image captioning based on the latent space. *Neurocomputing*, 546:126319.
- Patrick Esser, Sumith Kulal, Andreas Blattmann, Rahim Entezari, Jonas Müller, Harry Saini, Yam Levi, Dominik Lorenz, Axel Sauer, Frederic Boesel, and 1 others. 2024. Scaling rectified flow transformers for high-resolution image synthesis. In *Forty-first International Conference on Machine Learning*.
- Ruth C Fong and Andrea Vedaldi. 2017. Interpretable explanations of black boxes by meaningful perturbation. In *Proceedings of the IEEE international conference on computer vision*, pages 3429–3437.
- Yossi Gandelsman, Alexei A Efros, and Jacob Steinhardt. Interpreting clip’s image representation via text-based decomposition. In *The Twelfth International Conference on Learning Representations*.
- Yossi Gandelsman, Alexei A. Efros, and Jacob Steinhardt. 2024. [Interpreting the second-order effects of neurons in clip](#). *Preprint*, arXiv:2406.04341.
- Siobhan Mackenzie Hall, Fernanda Gonçalves Abrantes, Hanwen Zhu, Grace Sodunke, Aleksandar Shtedritski, and Hannah Rose Kirk. 2023. Visogender: A dataset for benchmarking gender bias in image-text pronoun resolution. *Advances in Neural Information Processing Systems*, 36:63687–63723.

- Lisa Anne Hendricks, Zeynep Akata, Marcus Rohrbach, Jeff Donahue, Bernt Schiele, and Trevor Darrell. 2016. Generating visual explanations. In *Computer Vision—ECCV 2016: 14th European Conference, Amsterdam, The Netherlands, October 11–14, 2016, Proceedings, Part IV 14*, pages 3–19. Springer.
- Lisa Anne Hendricks and Aida Nematzadeh. 2021. Probing image-language transformers for verb understanding. *arXiv preprint arXiv:2106.09141*.
- Dan Hendrycks, Steven Basart, Norman Mu, Saurav Kadavath, Frank Wang, Evan Dorundo, Rahul Desai, Tyler Zhu, Samyak Parajuli, Mike Guo, and 1 others. 2021a. The many faces of robustness: A critical analysis of out-of-distribution generalization. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 8340–8349.
- Dan Hendrycks, Kevin Zhao, Steven Basart, Jacob Steinhardt, and Dawn Song. 2021b. Natural adversarial examples. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 15262–15271.
- Phillip Howard, Avinash Madasu, Tiep Le, Gustavo Lujan Moreno, Anahita Bhiwandiwalla, and Vasudev Lal. 2024. Socialcounterfactuals: Probing and mitigating intersectional social biases in vision-language models with counterfactual examples. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 11975–11985.
- Kimmo Karkkainen and Jungseock Joo. 2021. Fairface: Face attribute dataset for balanced race, gender, and age for bias measurement and mitigation. In *Proceedings of the IEEE/CVF winter conference on applications of computer vision*, pages 1548–1558.
- Alexander Kirillov, Eric Mintun, Nikhila Ravi, Hanzi Mao, Chloe Rolland, Laura Gustafson, Tete Xiao, Spencer Whitehead, Alexander C Berg, Wan-Yen Lo, and 1 others. 2023. Segment anything. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4015–4026.
- Alex Krizhevsky, Geoffrey Hinton, and 1 others. 2009. Learning multiple layers of features from tiny images.
- Adam Dahlgren Lindström, Suna Bensch, Johanna Björklund, and Frank Drewes. 2021. Probing multimodal embeddings for linguistic properties: the visual-semantic case. *arXiv preprint arXiv:2102.11115*.
- Haotian Liu, Chunyuan Li, Yuheng Li, and Yong Jae Lee. 2024. Improved baselines with visual instruction tuning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 26296–26306.
- Ming Liu and Wensheng Zhang. 2025. [Is your video language model a reliable judge?](#) In *The Thirteenth International Conference on Learning Representations*.
- Huaishao Luo, Lei Ji, Ming Zhong, Yang Chen, Wen Lei, Nan Duan, and Tianrui Li. 2022. Clip4clip: An empirical study of clip for end to end video clip retrieval and captioning. *Neurocomputing*, 508:293–304.
- Avinash Madasu and Vasudev Lal. 2023. Is multimodal vision supervision beneficial to language? In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 2637–2642.
- Evelyn Mannix and Howard Bondell. 2024. Scalable and robust transformer decoders for interpretable image classification with foundation models. *arXiv preprint arXiv:2403.04125*.
- Prasanna Mayilvahanan, Roland S Zimmermann, Thaddäus Wiedemer, Evgenia Rusak, Attila Juhos, Matthias Bethge, and Wieland Brendel. In search of forgotten domain generalization. In *The Thirteenth International Conference on Learning Representations*.
- Lucas Moeller, Pascal Tilli, Ngoc Thang Vu, and Sebastian Pado. 2024. Explaining caption-image interactions in clip models with second-order attributions. *arXiv preprint arXiv:2408.14153*.
- Omkar M Parkhi, Andrea Vedaldi, Andrew Zisserman, and CV Jawahar. 2012. Cats and dogs. In *2012 IEEE conference on computer vision and pattern recognition*, pages 3498–3505. IEEE.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, and 1 others. 2021. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR.
- Christoph Schuhmann, Romain Beaumont, Richard Vencu, Cade Gordon, Ross Wightman, Mehdi Cherti, Theo Coombes, Aarush Katta, Clayton Mullis, Mitchell Wortsman, and 1 others. 2022. Laion-5b: An open large-scale dataset for training next generation image-text models. *Advances in neural information processing systems*, 35:25278–25294.
- Ramprasaath R Selvaraju, Michael Cogswell, Abhishek Das, Ramakrishna Vedantam, Devi Parikh, and Dhruv Batra. 2017. Grad-cam: Visual explanations from deep networks via gradient-based localization. In *Proceedings of the IEEE international conference on computer vision*, pages 618–626.
- Ashish Seth, Mayur Hemani, and Chirag Agarwal. 2023. Dear: Debiasing vision-language models with additive residuals. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6820–6829.
- Yang Shu, Xingzhuo Guo, Jialong Wu, Ximei Wang, Jianmin Wang, and Mingsheng Long. 2023. Clipood: Generalizing clip to out-of-distributions. In *International conference on machine learning*, pages 31716–31731. PMLR.

- K Simonyan, A Vedaldi, and A Zisserman. 2014. Deep inside convolutional networks: visualising image classification models and saliency maps. In *Proceedings of the International Conference on Learning Representations (ICLR)*. ICLR.
- Gabriela Ben Melech Stan, Raanan Yehezkel Rohekar, Yaniv Gurwicz, Matthew Lyle Olson, Anahita Bhiwandiwalla, Estelle Aflalo, Chenfei Wu, Nan Duan, Shao-Yen Tseng, and Vasudev Lal. 2024. Lvlm-intrepret: An interpretability tool for large vision-language models. *arXiv preprint arXiv:2404.03118*.
- Jun Xu, Tao Mei, Ting Yao, and Yong Rui. 2016. Msr-vtt: A large video description dataset for bridging video and language. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 5288–5296.
- Mengqi Xue, Qihan Huang, Haoifei Zhang, Lechao Cheng, Jie Song, Minghui Wu, and Mingli Song. 2022. Protopformer: Concentrating on prototypical parts in vision transformers for interpretable image recognition. *arXiv preprint arXiv:2208.10431*.
- Matthew D Zeiler and Rob Fergus. 2014. Visualizing and understanding convolutional networks. In *Computer Vision—ECCV 2014: 13th European Conference, Zurich, Switzerland, September 6–12, 2014, Proceedings, Part I 13*, pages 818–833. Springer.

A Concept Consistency Scores (CCS) for CLIP models.

We measure $CCS@K$ for all values of K i.e. $K \in [0, 5]$. Table 10 presents the Concept Consistency Score (CCS) distribution across various CLIP models, categorized by architecture size, patch size, and pre-training data. Several noteworthy trends emerge from this analysis. First, models pre-trained on larger and more diverse datasets (e.g., OpenCLIP-LAION2B) tend to exhibit a higher proportion of heads with $CCS@5$, indicating that a greater number of transformer heads are aligned with semantically meaningful concepts. For instance, the ViT-L-14 model trained on LAION2B shows the highest $CCS@5$ score of 0.328, suggesting that approximately 32.8% of heads are consistently associated with a single concept, reflecting strong concept alignment in these models.

Second, smaller models such as ViT-B-32 trained on OpenAI-400M demonstrate a significantly lower $CCS@5$ score (0.167) and a higher proportion of heads with lower CCS values (e.g., $CCS@0 = 0.021$), indicating weaker alignment of heads to consistent concepts. This observation implies that larger models with richer pre-training data are better at learning concept-specific representations, a key requirement for robust and interpretable multimodal reasoning.

Interestingly, when comparing models with the same architecture but different pre-training corpora, such as ViT-B-32 (OpenAI-400M vs. OpenCLIP-datacomp), we observe a higher $CCS@5$ score for datacomp (0.229) than OpenAI-400M (0.167), suggesting that dataset composition significantly affects the emergence of interpretable heads.

Moreover, progressive increases in CCS from $CCS@0$ to $CCS@5$ show how concept alignment varies within each model. For instance, while ViT-L-14 (OpenCLIP-LAION2B) has a low $CCS@0$ of 0.016, it steadily increases to a high $CCS@5$ of 0.328, suggesting that although a few heads are poorly aligned, a substantial fraction are highly consistent in capturing specific concepts.

In summary, these results demonstrate that the CCS metric effectively captures differences in conceptual alignment across models of varying size and pre-training datasets. Models with larger capacities and richer pre-training datasets tend to exhibit higher concept consistency, offering better interpretability and potentially stronger generalization abilities. This analysis underscores the value of

CCS as a diagnostic tool for evaluating and comparing the internal conceptual representations learned by CLIP-like models.

B Redundancy and Functional Duplication Analysis.

Previously in section 4.1, we observed that larger models degrade less when high-CCS (Concept Contribution Score) attention heads are pruned. This raises the question of whether larger models exhibit greater redundancy—i.e., whether multiple heads are attending to the same concepts, providing robustness against pruning.

To quantify this, we introduce the Concentrated Concept Ratio (CCR), a metric that captures the extent to which concepts are redundantly represented across multiple high-CCS heads.

$$CCR = \frac{C_{multi}}{H_{high}} \quad (1)$$

where H_{high} be the set of attention heads with high CCS, C_{multi} be the number of unique concepts that are concentrated on more than one head within H_{high} . A higher CCR implies greater redundancy—i.e., the same concepts are captured by multiple heads. Conversely, a lower CCR suggests a more distributed and unique representation of concepts across heads.

Table 11 presents CCR values across models of varying sizes. Interestingly, larger models (e.g., ViT-L variants) show lower CCR, indicating that they tend to distribute conceptual information more evenly across heads. This aligns with the empirical finding that these models are more resilient to pruning of high-CCS heads.

To further examine whether redundancy exists across layers (functional duplication), we perform a layer-wise CCR analysis restricted to the last four transformer layers of each model. This focus is motivated by prior work (Gandelsman et al.), which shows that only the final four layers significantly affect model outputs. CCR is computed independently for each of the last four layers. However, the analysis reveals no consistent trend indicating that deeper layers exhibit greater concept duplication. This suggests that redundancy across layers is not a dominant factor and that the observed robustness in larger models is more likely due to distributed head-level representations rather than functional duplication across layers.

Model	Model size	Patch size	Pre-training data	CCS@0	CCS@1	CCS@2	CCS@3	CCS@4	CCS@5
CLIP	B	32	OpenAI-400M	0.021	0.062	0.167	0.271	0.312	0.167
CLIP	B	32	OpenCLIP-datacomp	0.104	0.062	0.208	0.189	0.208	0.229
CLIP	B	16	OpenAI-400M	0.021	0.062	0.125	0.292	0.292	0.208
CLIP	B	16	OpenCLIP-LAION2B	0.062	0.062	0.105	0.25	0.25	0.271
CLIP	L	14	OpenAI-400M	0.062	0.109	0.172	0.204	0.203	0.25
CLIP	L	14	OpenCLIP-LAION2B	0.016	0.031	0.109	0.219	0.297	0.328

Table 10: **Concept Consistency Score (CCS) for CLIP models.**

Model	Concentrated Concept Ratio (CCR)
ViT-B-32-OpenAI	0.250
ViT-B-32-datacomp	0.273
ViT-B-16-OpenAI	0.200
ViT-B-16-LAION	0.231
ViT-L-14-OpenAI	0.125
ViT-B-14-LAION	0.143

Table 11: **Concentrated Concept Ratio (CCR) across different CLIP model variants.** Lower CCR in larger models (e.g., ViT-L-14) suggests more distributed concept representations.

ViT-B-32-OpenAI L8.H11 (Descriptive), L9.H2 (Objects), L9.H3 (Descriptions), L10.H8 (Locations), L11.H1 (Objects), L11.H5 (Colors), L11.H7 (Objects), L11.H9 (Locations)
ViT-B-32-datacomp L8.H1 (Objects), L8.H3 (Subjects), L8.H10 (Objects), L9.H3 (Subjects), L9.H10 (Objects), L10.H7 (Locations), L10.H11 (Objects), L11.H3 (Colors), L11.H4 (Colors), L11.H9 (Colors), L11.H10 (Objects)
ViT-B-16-OpenAI L8.H5 (Visual), L8.H8 (Visual), L10.H5 (Subjects), L10.H7 (Settings), L11.H0 (Creative), L11.H3 (Settings), L11.H4 (Stylistic), L11.H6 (Locations), L11.H7 (Colors), L11.H11 (Animals)
ViT-B-16-LAION L8.H6 (Descriptions), L8.H7 (Descriptions), L9.H0 (Themes), L9.H1 (Aesthetics), L9.H3 (Descriptive), L10.H5 (Artwork), L10.H10 (Locations), L11.H0 (Locations), L11.H2 (Descriptions), L11.H6 (Locations), L11.H7 (Objects), L11.H8 (Objects), L11.H10 (Colors)
ViT-L-14-OpenAI L20.H2 (Locations), L20.H12 (Descriptions), L21.H0 (Locations), L21.H1 (Locations), L21.H8 (Expressions), L21.H13 (Locations), L21.H15 (Locations), L22.H1 (Objects), L22.H2 (Locations), L22.H5 (Locations), L22.H9 (Subjects), L22.H13 (Animals), L22.H15 (Locations), L23.H4 (Objects), L23.H10 (Locations), L23.H11 (Colors)
ViT-L-14-LAION L20.H4 (Subjects), L20.H14 (Descriptions), L21.H0 (Colors), L21.H1 (Locations), L21.H5 (Descriptive), L21.H9 (Colors), L21.H11 (Locations), L22.H0 (Patterns), L22.H1 (Shapes), L22.H3 (Objects), L22.H5 (Visual), L22.H6 (Animals), L22.H8 (Letters), L22.H10 (Colors), L22.H12 (Landscapes), L22.H13 (Locations), L23.H4 (People), L23.H5 (Nature), L23.H6 (Locations), L23.H8 (Colors), L23.H9 (Descriptive)

Table 12: **Full List of high-CCS heads of all CLIP models.**

ViT-B-32-OpenAI

L8.H1 (Artistic), L8.H2 (Objects), L8.H6 (Photography), L8.H9 (Styles), L8.H10 (Perspective), L9.H1 (Subjects), L9.H11 (Settings), L10.H0 (Objects), L10.H3 (Locations), L10.H7 (Locations), L11.H6 (Descriptions), L11.H10 (Locations), L11.H11 (Locations)

ViT-B-32-datacamp

L8.H0 (Environments), L8.H7 (Creativity), L9.H6 (Colors), L10.H5 (Art), L10.H6 (Descriptions), L10.H8 (Locations), L10.H9 (Descriptions), L11.H2 (Subjects), L11.H8 (Qualities)

ViT-B-16-OpenAI

L8.H1 (Artistic), L8.H2 (Photography), L8.H4 (Styles), L8.H6 (Artwork), L8.H7 (Photography), L8.H9 (Light), L9.H4 (Photography), L9.H6 (Artforms), L9.H10 (Elements), L10.H3 (Locations), L10.H8 (Colors), L10.H9 (Artwork), L11.H5 (Objects), L11.H8 (Effects)

ViT-B-16-LAION

L8.H0 (Locations), L8.H8 (Text), L8.H9 (Photography), L9.H7 (Artistic), L9.H8 (Settings), L9.H11 (Descriptions), L10.H2 (Nature), L10.H3 (Location), L10.H7 (Expressions), L11.H3 (Settings), L11.H9 (Numbers), L11.H11 (Letters)

ViT-L-14-OpenAI

L20.H0 (Locations), L20.H3 (Locations), L20.H7 (Communication), L20.H8 (Vehicles), L20.H10 (Locations), L21.H4 (Photography), L21.H6 (People), L21.H10 (Locations), L22.H3 (Countries), L22.H12 (Professions), L23.H3 (Patterns), L23.H9 (Creativity), L23.H15 (Visual)

ViT-L-14-LAION

L20.H0 (Locations), L20.H1 (Locations), L20.H2 (Locations), L20.H8 (Locations), L20.H9 (Locations), L20.H11 (Aesthetics), L20.H15 (Descriptions), L21.H12 (Photography), L21.H14 (Locations), L22.H9 (Activities), L22.H14 (Colors), L22.H15 (Emotions), L23.H0 (Materials), L23.H3 (Settings)

Table 13: **Full List of medium-CCS heads of all CLIP models.**

ViT-B-32-OpenAI

L8.H5 (Patterns), L9.H9 (Ambiance), L11.H0 (Diverse), L11.H8 (Word)

ViT-B-32-datacamp

L8.H2 (Images), L8.H4 (Varied), L8.H9 (Varied), L9.H4 (Variety), L9.H5 (Professions), L11.H0 (Diverse), L11.H1 (Varied), L11.H11 (Settings)

ViT-B-16-OpenAI

L8.H0 (Diversity), L9.H3 (Locations), L10.H6 (Body parts), L11.H2 (Perspective)

ViT-B-16-LAION

L8.H4 (Variety), L8.H5 (Varied), L8.H10 (Diverse), L9.H2 (Textures), L10.H6 (Photography), L10.H8 (Traits)

ViT-L-14-OpenAI

L20.H1 (Diverse), L20.H4 (Diversity), L20.H6 (Items), L20.H15 (Diverse), L21.H2 (Diversity), L21.H3 (Diverse), L22.H0 (Occupations), L22.H4 (Settings), L22.H6 (Weather), L22.H14 (Items), L23.H5 (Diversity)

ViT-L-14-LAION

L20.H13 (Photography), L21.H6 (Professions), L23.H1 (Diverse)

Table 14: **Full List of low-CCS heads of all CLIP models.**

Occupations
biologist, composer, economist, mathematician, model, poet, reporter, zoologist, artist, coach, athlete, audiologist, judge, musician, therapist, banker, ceo, consultant, prisoner, assistant, boxer, commander, librarian, nutritionist, realtor, supervisor, architect, priest, guard, magician, producer, teacher, lawyer, paramedic, researcher, physicist, pediatrician, surveyor, laborer, statistician, dietitian, sailor, tailor, attorney, army, manager, baker, recruiter, clerk, entrepreneur, sheriff, policeman, businessperson, chief, scientist, carpenter, florist, optician, salesperson, umpire, painter, guitarist, broker, pensioner, soldier, astronaut, dj, driver, engineer, cleaner, cook, housekeeper, swimmer, janitor, pilot, mover, handyman, firefighter, accountant, physician, farmer, bricklayer, photographer, surgeon, dentist, pianist, hairdresser, receptionist, waiter, butcher, videographer, cashier, technician, chemist, blacksmith, dancer, doctor, nurse, mechanic, chef, plumber, bartender, pharmacist, electrician

Table 15: **Full list of occupations used for evaluating biases on FairFace and SocialCounterFactuals datasets.**

Prompt	Example
A <occupation>	A biologist
A photo of <occupation>	A photo of biologist
A picture of <occupation>	A picture of biologist
An image of <occupation>	An image of biologist

Table 16: **Prompts used for measuring biases on FairFace and SocialCounterFactuals datasets.**