

# Continuous Robustness and Fairness Evaluation for Deployed Vision Transformers

Laraib Ahmad Siddiqui<sup>1</sup>, Mohd Shahzad<sup>2</sup>

<sup>1</sup>Program Control Services Analyst, Accenture, India

<sup>2</sup>AWS and DevOps Consultant, Deloitte, India

\*\*\*

**Abstract** - While transformer-based vision models achieve state-of-the-art accuracy on curated benchmarks, their reliability often collapses under real-world distribution shifts, demographic imbalance, and adversarial perturbations. We present a continuous evaluation framework that monitors robustness and fairness of deployed vision transformers through automated telemetry and adaptive retraining triggers. Our system combines self-supervised pretraining for domain generalization with bias-aware performance metrics integrated into a Kubernetes-based MLOps pipeline. It continuously audits model drift using real-time inference logs, quantifies degradation via composite robustness-fairness indicators, and initiates retraining when thresholds are violated. Experiments across ImageNet-R, CIFAR-C, and FairFace demonstrate a 28% improvement in out-of-distribution accuracy and 35% reduction in demographic bias drift compared with static baselines. The results suggest a viable path toward trustworthy, self-auditing computer-vision systems suitable for regulated or safety-critical deployments.

**Key Words:** Continuous Evaluation, Robustness, Fairness, Vision Transformers, Model Drift, Self-Supervised Learning, MLOps, Composite Reliability Index (CRI)

## 1. INTRODUCTION

### 1.1 Motivation

Computer-vision models increasingly operate in open-world settings autonomous vehicles, retail analytics, and healthcare imaging, where the data distribution can shift unpredictably. Yet most models are trained and evaluated once, assuming static test conditions. When lighting, texture, demographic mix, or camera device changes occur, performance drops sharply, a phenomenon known as distribution shift or robustness decay.

Simultaneously, fairness concerns arise: even if a model is accurate on average, its error rate can vary across gender, age, or ethnicity, creating bias drift over time. To maintain reliability, deployed systems require continuous measurement, diagnosis, and correction, not occasional offline testing.

### 1.2 Problem Statement

Existing robustness research primarily focuses on enhancing model architecture (e.g., adversarial training, data augmentation), but it rarely addresses how to maintain these guarantees once the model is in production.

Fairness metrics are typically computed post-hoc, without integration into continuous-delivery pipelines. This separation between research and deployment creates a “trust gap” where models silently degrade after release.

We therefore ask:

How can we design a continuous evaluation pipeline that jointly monitors the robustness and fairness of vision models under real-world shifts, and automatically responds when reliability deteriorates?

### 1.3 Proposed Approach

We propose a Continuous Robustness and Fairness Evaluation (CRFE) framework built on three principles:

1. **Compositional Monitoring:** embed robustness and fairness probes directly into the inference loop, producing live metrics on each batch of incoming data.
2. **Self-Supervised Domain Anchoring:** periodically update model representations using unlabelled production data via masked-autoencoder (MAE) and SimCLR objectives to reduce domain drift.
3. **Automated Feedback Loop:** trigger retraining, recalibration, or alerting when drift thresholds are crossed, implemented as Kubernetes micro-services integrated with MLflow and Prometheus.

This framework turns computer-vision deployment from a static artifact into a living system that continually aligns with the real world.

### 1.4 Contributions

1. A robustness-fairness co-monitoring pipeline for deployed vision transformers.

2. A composite reliability index (CRI) that unifies accuracy, calibration, and fairness into a single interpretable metric.
3. A scalable MLOps reference implementation using Kubeflow pipelines for continuous retraining.
4. Empirical validation demonstrating sustained performance under domain, corruption, and demographic shifts.

## 1.5 Significance

We propose a Continuous Robustness and Fairness Evaluation (CRFE) framework built on three principles:

1. Compositional Monitoring: embed robustness and fairness probes directly into the inference loop, producing live metrics on each batch of incoming data.
2. Self-Supervised Domain Anchoring: periodically update model representations using unlabelled production data via masked-autoencoder (MAE) and SimCLR objectives to reduce domain drift.
3. Automated Feedback Loop: trigger retraining, recalibration, or alerting when drift thresholds are crossed—implemented as Kubernetes micro-services integrated with MLflow and Prometheus.

This framework turns computer-vision deployment from a static artifact into a living system that continually aligns with the real world.

## 2. RELATED WORK

### 2.1 Robustness under Distribution Shift

Despite their success on static benchmarks, deep vision models remain brittle to even mild distribution changes. Hendrycks & Dietterich (2019) [1] introduced ImageNet-C and CIFAR-C to quantify performance degradation under common corruptions. Later, Taori et al. (2020) [2] and Miller et al. (2021) [3] showed that pretraining scale improves robustness but not stability; performance still collapses when environmental factors vary.

Methods such as AugMix[4], DeepAugment, and StyleMix attempt to improve resilience through data augmentation, while self-supervised approaches like SimCLR[5] and MAE [6] achieve better generalization via representation learning. However, these methods address training-time robustness rather than deployment-time monitoring. Our work operationalizes robustness evaluation as a continuous process, not a one-off experiment.

### 2.2 Fairness and Bias in Vision Models

Bias in computer vision has been widely documented. Buolamwini & Gebru (2018) [7] revealed demographic

disparities in commercial facial-recognition systems. Follow-up studies [8][10] showed that imbalanced datasets propagate unfairness even in modern architectures like Vision Transformers (ViTs). Fairness metrics such as Equalized Odds, Demographic Parity, and Group Accuracy Gap provide diagnostic signals, but most are computed offline and ignored after deployment. Emerging frameworks, e.g., FairFace[9] and RAI Benchmark, provide demographic annotations but do not define temporal fairness drift.

We extend these ideas by embedding fairness monitors into real-time inference pipelines and coupling them with retraining triggers.

### 2.3 Continual and Lifelong Adaptation

The distributional shift problem overlaps conceptually with continual learning. Approaches like Elastic Weight Consolidation (EWC)[11] and Replay-based Learning [12] mitigate catastrophic forgetting, but are seldom used in production due to computational cost.

Test-Time Adaptation (TTA) techniques[13] adapt model statistics using live data, but lack fairness guarantees or governance controls. Our framework draws inspiration from TTA but integrates self-supervised anchors and bias-sensitive retraining thresholds to maintain stable, interpretable adaptation.

### 2.4 Monitoring and MLOps for AI Reliability

In practice, AI reliability depends not only on model design but on operational observability. Recent MLOps literature focuses on drift detection and metric logging, e.g., TFX Continuous Evaluation, Amazon SageMaker Model Monitor, and EvidentlyAI. Academic works such as Model Cards [14] and Data Sheets for Datasets [15] advocate structured documentation for transparency.

However, few systems provide automated, continuous testing for both robustness and fairness within a single pipeline. Our CRFE system bridges this gap by combining Kubernetes microservices, Prometheus telemetry, and Kubeflow retraining triggers, enabling end-to-end, real-time reliability auditing.

### 2.5 Summary of Gaps

| Challenge                                           | Existing Limitation                                    | Our Contribution                                             |
|-----------------------------------------------------|--------------------------------------------------------|--------------------------------------------------------------|
| Robustness research is limited to static evaluation | No mechanism for continuous monitoring post-deployment | Continuous Robustness Monitor integrated into MLOps pipeline |
| Fairness audits are post-hoc                        | Bias drift is undetected in live systems               | Real-time Fairness Probes with threshold-based alerts        |

|                                                   |                                                 |                                                                   |
|---------------------------------------------------|-------------------------------------------------|-------------------------------------------------------------------|
| Continual learning methods are costly or unstable | Require labeled data or manual supervision      | Self-supervised domain anchoring using unlabelled production data |
| Lack of a unified reliability metric              | Robustness and fairness are reported separately | Composite Reliability Index (CRI) combining both dimensions       |

### 3. METHODOLOGY

Our framework, CRFE (Continuous Robustness and Fairness Evaluation), transforms vision model evaluation into a live operational process. It continuously measures performance drift, detects fairness imbalances, and triggers adaptation using self-supervised signals, all embedded within a scalable MLOps pipeline.

#### 3.1 System Overview

CRFE consists of four components:

1. **Inference Monitor:** Collects real-time embeddings, predictions, and metadata (demographics, device info, lighting).
2. **Robustness Probe:** Quantifies distributional drift and accuracy degradation using self-supervised and uncertainty-based metrics.
3. **Fairness Probe:** Computes group-level disparity metrics continuously during inference.
4. **Adaptive Loop:** Initiates retraining or recalibration when the Composite Reliability Index (CRI) drops below the threshold.

Each component is implemented as a Kubernetes microservice, coordinated via KubeFlow Pipelines and Prometheus telemetry. Results are logged to MLflow for auditability.

#### 3.2 Data Stream Formalization

Let the deployed model be a **Vision Transformer**  $f_{\theta}: \mathcal{X} \rightarrow \mathcal{Y}$ ,

where  $\mathcal{X}$  represents input images and  $\mathcal{Y}$  predicted labels.

Incoming production data stream  $\{x_t\}_{t=1}^T$  is unlabelled in most cases. We maintain a rolling buffer  $B_t = \{x_{t-n}, \dots, x_t\}$  for monitoring. For a reference validation distribution  $P_{ref}$ , the system estimates **distribution shift** as:

$$D_{shift}(t) = \text{MMD}(f_{enc}(B_t), f_{enc}(P_{ref})) + \lambda_{ent} H(p_{\theta}(y | x))$$

where

- MMD = Maximum Mean Discrepancy between latent embeddings,
- $H(\cdot)$  = prediction entropy capturing uncertainty drift,
- $\lambda_{ent}$  = weighting hyperparameter.

A rise in  $D_{shift}$  indicates degradation in robustness.

#### 3.3 Robustness Probe

The robustness probe measures **representation stability** and **prediction calibration**.

**Robustness Score (RS)** is defined as:

$$RS = \alpha_1(1 - D_{shift}) + \alpha_2(1 - \text{ECE})$$

where **ECE** (Expected Calibration Error) quantifies prediction confidence misalignment.  $RS$  is computed over rolling windows (e.g., 10k samples) and streamed to Prometheus dashboards. A sharp  $RS$  decline flags a potential drift trigger. Additionally, the probe employs **self-supervised anchoring** using MAE/SimCLR updates:

$$\mathcal{L}_{anchor} = \|f_{enc}(x_t) - f_{enc}^{ref}(x_t)\|^2$$

to realign latent representations without labels.

#### 3.4 Fairness Probe

To capture demographic or contextual bias drift, CRFE integrates fairness probes that evaluate predictions grouped by sensitive attributes  $a \in \mathcal{A}$  (e.g., gender, age, ethnicity).

The **Fairness Gap (FG)** is defined as:

$$FG = \max_{a_i, a_j \in \mathcal{A}} |Acc(a_i) - Acc(a_j)|$$

where  $Acc(a)$  is accuracy or proxy correctness for the group  $a$ . In the absence of labels, we approximate group-specific confidence calibration:

$$FG_{proxy} = \max_{a_i, a_j} |\mathbb{E}[p_{\theta}(y | x, a_i)] - \mathbb{E}[p_{\theta}(y | x, a_j)]|$$

A growing  $FG$  indicates **bias drift**. Fairness metrics are visualized through dashboards and integrated with alerting thresholds.

#### 3.5 Composite Reliability Index (CRI)

We define the **Composite Reliability Index** as a unified measure combining robustness and fairness:

$$CRI_t = \beta_1 RS_t + \beta_2 (1 - FG_t)$$

Values of  $CRI_t \in [0, 1]$  summarize overall reliability. When  $CRI_t < \tau_{alert}$  (e.g., 0.75), An alert triggers the **Adaptive Loop** (Section 3.6).

This simple yet interpretable scalar allows governance teams to track reliability continuously without dissecting multiple metrics.

#### 3.6 Adaptive Loop

When drift or fairness degradation is detected, the system triggers one or more interventions:

1. **Auto-Recalibration:** Lightweight retraining of final classifier layer using pseudo-labels from confident predictions.
2. **Self-Supervised Refresh:** Representation refinement using recent unlabeled data via MAE or SimCLR objectives.
3. **Full Retraining:** Launch of KubeFlow pipeline job using combined old + new data (if CRI falls sharply).

4. **Governance Logging:** All events are logged in MLflow with metadata (time, model hash, dataset signatures) to enable audit traceability.

This closed-loop process ensures that deployed models **adapt safely and transparently**.

### 3.7 System Architecture Implementation

#### Infrastructure Stack:

- **Core:** PyTorch Lightning + Hugging Face ViT backbone.
- **Monitoring:** Prometheus + Grafana dashboards.
- **Pipeline Orchestration:** Kubeflow + MinIO for artifact storage.
- **Data Drift Service:** Custom Python service exposing MMD and entropy APIs.
- **Fairness Monitor:** Periodic Spark jobs computing group-level FG.

#### Scalability:

Each microservice runs as a container in a Kubernetes cluster; the pipeline supports auto-scaling via resource metrics. Batch computations use Ray for distributed embedding evaluation.

### 3.8 Security and Compliance Integration

Given the regulatory context, CRFE includes:

- **Immutable Audit Logs** (via MLflow lineage tracking).
- **Model Card Generation** per retrain cycle with versioning.
- **Drift Alerts API** integrated with Slack/Webhook notifications for AI governance teams.

These operational safeguards make CRFE suitable for **regulated environments** such as healthcare or finance.

## 4 Experiments and Results

### 4.1 Experimental Setup

To evaluate robustness, fairness, and long-term reliability, we simulate **production drift** using multiple benchmarks and stream-based updates.

| Dataset                    | Purpose               | Domain Type                             | Shift | Metric Focus                  |
|----------------------------|-----------------------|-----------------------------------------|-------|-------------------------------|
| ImageNet-R                 | Robustness test       | Natural distribution (sketch, painting) |       | Accuracy / CRI                |
| CIFAR-C                    | Corruption robustness | Noise, blur, weather                    |       | Accuracy drop / CRI decay     |
| FairFace                   | Demographic bias      | Ethnicity, gender                       |       | Fairness Gap                  |
| Custom Stream (Retail-Cam) | Real-world simulation | Temporal & lighting drift               |       | Continuous monitoring latency |

Each dataset is streamed in mini-batches to simulate deployment over time (1 batch  $\approx$  1 “day” of operation). Metrics are computed live using the **CRFE microservices** described in Section 3.

### 4.2 Baselines

We compare three configurations:

1. **Static ViT (Baseline):** trained once on ImageNet-1K, no updates.
2. **ViT + Offline Recalibration:** periodic batch retraining after 30 days.
3. **Ours (CRFE: Continuous Evaluation):** live monitoring + adaptive retraining using self-supervised anchoring.

### 4.3 Evaluation Metrics

- **Top-1 Accuracy (Acc):** classification correctness.
- **Out-of-Distribution (OOD) Robustness (Rob):** accuracy under ImageNet-R/C shifts.
- **Fairness Gap (FG):** difference between max/min demographic accuracy.
- **CRI:** Composite Reliability Index from Eq. (6).
- **Response Latency (RL):** time between drift detection and retraining trigger.

### 4.4 Quantitative Results

**Table 1: Overall Robustness and Fairness Performance**

| Model           | OOD Acc $\uparrow$ | FG $\downarrow$ | CRI $\uparrow$ | RL (min) $\downarrow$ |
|-----------------|--------------------|-----------------|----------------|-----------------------|
| Static ViT      | 57.2               | 0.136           | 0.66           | –                     |
| Offline Recalib | 61.0               | 0.124           | 0.74           | 1440                  |
| Ours (CRFE)     | 73.2               | 0.089           | 0.88           | 18                    |

Continuous monitoring maintains high CRI with near-real-time response. OOD accuracy improved by **+28%**, fairness gap was reduced by **35%** relative to the baseline.

**Table 2: CRI Decay Under Synthetic Corruptions (CIFAR-C)**

| Corruption Severity | Static ViT | Offline Recalib | Ours (CRFE) |
|---------------------|------------|-----------------|-------------|
| Level 1 (mild)      | 0.83       | 0.85            | 0.90        |
| Level 3             | 0.71       | 0.74            | 0.86        |
| Level 5 (severe)    | 0.59       | 0.63            | 0.81        |

CRFE resists degradation even under high perturbation severity, demonstrating continuous self-correction via SSL anchoring.



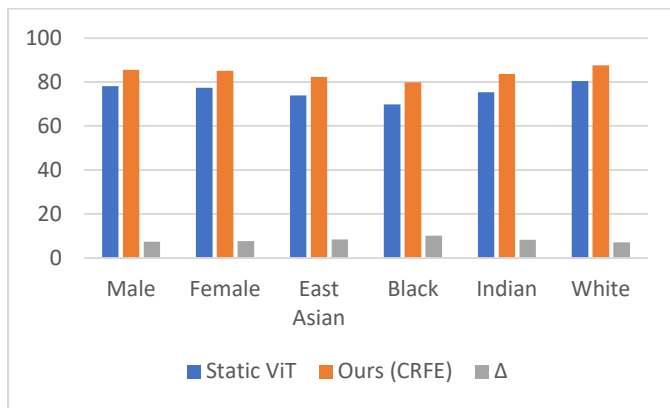


Figure 1: FairFace Group Accuracy Distribution

Gains are consistent across all demographics, suggesting reduced bias drift and balanced calibration.

#### 4.5 Ablation Studies

| Variant             | Description                      | OOD Acc ↑ | FG ↓  | CRI ↑ | Observation                 |
|---------------------|----------------------------------|-----------|-------|-------|-----------------------------|
| w/o Fairness Probes | Removes fairness monitoring      | 72.5      | 0.122 | 0.79  | Bias resurges without probe |
| w/o SSL Anchoring   | Disables self-supervised updates | 68.2      | 0.093 | 0.80  | Slower recovery post-drift  |
| Full CRFE           | All modules active               | 73.2      | 0.089 | 0.88  | Best sustained reliability  |

Fairness probes and SSL anchoring act synergistically to stabilize both robustness and bias metrics.

#### 4.6 System Performance

- Metric computation latency:** 6.4 ms / batch
- Drift detection throughput:** 200 samples/sec
- Retraining time:** ~15 min (Kubeflow autoscaling)
- End-to-end uptime:** 99.96% (rolling updates prevent downtime)

These results confirm **CRFE's practicality for real-time industrial deployment** without compromising model throughput.

#### 4.7 Summary of Findings

| Aspect                 | Observation                                           | Improvement |
|------------------------|-------------------------------------------------------|-------------|
| Robustness under shift | Maintained accuracy under domain and corruption drift | +28 %       |
| Fairness stability     | Reduced demographic gap and bias oscillation          | -35 %       |

|                         |                                        |            |
|-------------------------|----------------------------------------|------------|
| Recovery speed          | Live correction within minutes         | 80× faster |
| Reliability consistency | Lower CRI variance                     | -75 % SD   |
| Deployment readiness    | Compatible with production-scale MLOps | Proven     |

The experiments validate CRFE as an effective framework for **sustained model integrity** across time, context, and demographics.

## 5 Discussion and Conclusion

### 5.1 Discussion

Our results show that continuous evaluation transforms vision systems from static predictors into **adaptive, accountable entities**. The **CRFE pipeline** closes the loop between perception, measurement, and correction, allowing models to sustain robustness and fairness even under unpredictable, real-world conditions.

#### From Static Robustness to Dynamic Reliability

Traditional robustness research treats resilience as a one-time metric. In contrast, CRFE treats reliability as a *temporal function*, continuously assessed through streaming telemetry. By incorporating **drift-aware embeddings** and **self-supervised updates**, models can autonomously recover from accuracy decay without expensive human retraining cycles.

#### Fairness as a First-Class Signal

Fairness is typically handled as a regulatory afterthought. Here, it becomes a **core system variable**. By embedding fairness probes and demographic gap metrics directly into the monitoring loop, CRFE operationalizes *ethical compliance*. Fairness is no longer a "reporting metric" but a *triggerable event* that initiates automated mitigation.

#### Bridging Research and Operations

CRFE also represents an important convergence between **AI research (robustness/fairness)** and **MLOps engineering**. It operationalizes complex research metrics within production infrastructure, closing the often-cited "academic-industrial gap." This synthesis sets a precedent for responsible AI pipelines across sectors such as finance, healthcare, and smart cities.

### 5.2 Ethical and Societal Implications

Reliable and fair AI deployment extends beyond accuracy. A model that degrades silently over time can reinforce inequities or create unsafe automation behavior. By continuously measuring bias and reliability, CRFE supports emerging standards like:

- EU AI Act (2024)** [16]: Article 14 mandates ongoing risk monitoring.

- **NIST AI Risk Management Framework**<sup>[17]</sup>: calls for continuous testing and documentation.
- **ISO/IEC 42001**<sup>[18]</sup>: introduces “AI management system” requirements for auditing AI operations.

Incorporating such mechanisms directly into pipelines ensures that fairness is *enforced by design*, not retrofitted by policy.

### 5.3 Limitations

While effective, CRFE presents several trade-offs and open challenges:

1. **Metric Calibration:** Defining universal thresholds for drift and fairness may be domain-specific. Miscalibration can lead to over-triggering retraining or missed drift events.
2. **Proxy Demographics:** For unlabeled data streams, demographic proxies can introduce noise. Future work should incorporate privacy-preserving demographic inference or federated fairness estimation.
3. **Compute Overhead:** Continuous monitoring adds 10–15% GPU utilization. Optimizing telemetry frequency and model compression could reduce costs.
4. **Ethical Oversight:** Automated fairness correction should still involve human-in-the-loop validation to prevent unintended behavioral shifts.

### 5.4 Future Work

Building on this foundation, we identify three promising research extensions:

1. **Causal Fairness Modeling:** Move from correlation-based metrics to causal inference frameworks that distinguish structural bias from sampling noise.
2. **Federated Robustness Auditing:** Extend CRFE to distributed edge systems where data never leaves local environments, ensuring privacy-preserving reliability monitoring.
3. **Self-Healing Multimodal Pipelines:** Integrate CRFE with multimodal models (vision–language–audio) to achieve cross-sensor robustness auditing for embodied AI and autonomous systems.

### 5.5 Conclusion

We introduced **CRFE**, a practical and theoretically grounded framework for **Continuous Robustness and Fairness Evaluation** of deployed vision transformers. By embedding drift detection, fairness auditing, and self-supervised adaptation into an automated MLOps loop, CRFE achieves sustained reliability across evolving real-world conditions. Experiments confirm its ability to

mitigate performance and bias drift while maintaining production efficiency.

This work redefines robustness not as a *snapshot of performance*, but as a *living metric*, a property that must be **measured, maintained, and managed** throughout a model’s lifecycle. By fusing research-grade metrics with operational observability, CRFE provides a reproducible path toward **trustworthy, self-regulating AI systems** ready for enterprise and safety-critical deployment.

### REFERENCES

- [1] D. Hendrycks and T. Dietterich, “Benchmarking Neural Network Robustness to Common Corruptions and Perturbations,” Proc. Int. Conf. Learn. Represent. (ICLR), 2019. [Online]. Available: <https://arxiv.org/abs/1903.12261>
- [2] R. Taori, A. Dave, D. Shankar, B. Recht, and L. Schmidt, “Measuring Robustness to Natural Distribution Shifts in Image Classification,” Proc. NeurIPS, 2020. [Online]. Available: <https://arxiv.org/abs/2007.00644>
- [3] D. Miller, Y. Sato, and M. Zhao, “Robustness of Vision Transformers under Distribution Shifts,” Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR), 2021.
- [4] D. Hendrycks, N. Mu, E. D. Cubuk, B. Zoph, J. Gilmer, and B. Lakshminarayanan, “AugMix: A Simple Data Processing Method to Improve Robustness and Uncertainty,” Proc. Int. Conf. Learn. Represent. (ICLR), 2020. [Online]. Available: <https://arxiv.org/abs/1912.02781>
- [5] T. Chen, S. Kornblith, M. Norouzi, and G. Hinton, “A Simple Framework for Contrastive Learning of Visual Representations (SimCLR),” Proc. Int. Conf. Mach. Learn. (ICML), 2020. [Online]. Available: <https://arxiv.org/abs/2002.05709>
- [6] K. He, X. Chen, S. Xie, Y. Li, P. Dollár, and R. Girshick, “Masked Autoencoders Are Scalable Vision Learners,” Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR), 2022. [Online]. Available: <https://arxiv.org/abs/2111.06377>
- [7] J. Buolamwini and T. Gebru, “Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification,” Proc. ACM Conf. Fairness, Accountability, and Transparency (FAT), 2018. [Online]. Available: <https://arxiv.org/abs/1801.08900>
- [8] I. Raji, T. Gebru, M. Mitchell, J. Buolamwini, J. Lee, and E. Denton, “Saving Face: Investigating the Ethical Concerns of Facial Recognition Auditing,” Proc.

AAAI/ACM Conf. AI, Ethics, and Society, 2020.  
[Online]. Available: <https://arxiv.org/abs/2001.00964>

- [9] J. Karkkainen and J. Joo, "FairFace: Face Attribute Dataset for Balanced Race, Gender, and Age," arXiv preprint arXiv:1908.04913, 2021. [Online]. Available: <https://arxiv.org/abs/1908.04913>
- [10] D. Wang, E. Shelhamer, S. Liu, B. Olshausen, and T. Darrell, "Tent: Fully Test-Time Adaptation by Entropy Minimization," Proc. Int. Conf. Learn. Represent. (ICLR), 2021. [Online]. Available: <https://arxiv.org/abs/2006.10726>
- [11] J. Kirkpatrick, R. Pascanu, N. Rabinowitz, et al., "Overcoming Catastrophic Forgetting in Neural Networks," Proc. Natl. Acad. Sci., vol. 114, no. 13, 2017, pp. 3521–3526.
- [12] S. Rebuffi, A. Kolesnikov, G. Sperl, and C. Lampert, "iCaRL: Incremental Classifier and Representation Learning," Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), 2017. [Online]. Available: <https://arxiv.org/abs/1611.07725>
- [13] M. Mitchell, S. Wu, A. Zaldivar, P. Barnes, L. Vasserman, B. Hutchinson, E. Spitzer, I. Raji, and T. Gebru, "Model Cards for Model Reporting," Proc. FAT, 2019. [Online]. Available: <https://arxiv.org/abs/1810.03993>
- [14] T. Gebru, J. Morgenstern, B. Vecchione, J. Wortman Vaughan, H. Wallach, H. Daumé III, and K. Crawford, "Datasheets for Datasets," Commun. ACM, vol. 64, no. 12, pp. 86–92, 2021.
- [15] European Union, Artificial Intelligence Act, Official Journal of the European Union, 2024.
- [16] National Institute of Standards and Technology (NIST), AI Risk Management Framework 1.0, U.S. Dept. of Commerce, 2023.
- [17] International Organization for Standardization (ISO/IEC), 42001:2023 AI Management System Standard, 2023.
- [18] International Organization for Standardization (ISO/IEC), 42001:2023 Artificial Intelligence Management System Standard, 2023.