



A Cognitive-Inspired Algorithm for Mitigating Temporal Bias in Artificial Intelligence-Generated Watch Images Using Diffusion Models

Noor F. Mohammed¹, Mohammed Safar^{1*}, Rawan A. AlRashid Agha²

¹ Information Techniques and Computer Networks Engineering Department, Technical Engineering College for Computer and AI-Kirkuk, Northern Technical University, Kirkuk 36001, Iraq

² Department of Computer Science and Engineering, School of Science and Engineering, University of Kurdistan Hewler (UKH), Erbil 44001, Iraq

Corresponding Author Email: mohammed.sefer@ntu.edu.iq

Copyright: ©2026 The authors. This article is published by IIETA and is licensed under the CC BY 4.0 license (<http://creativecommons.org/licenses/by/4.0/>).

<https://doi.org/10.18280/isi.310626>

ABSTRACT

Received: 18 October 2025

Revised: 16 December 2025

Accepted: 1 January 2026

Available online: 30 June 2026

Keywords:

generative artificial intelligence, diffusion models, artificial intelligence - generated image bias, temporal bias mitigation, cognitive-inspired computing, stable diffusion, image generation

Generative artificial intelligence (AI) systems have achieved remarkable progress in image synthesis; however, they often inherit visual biases from training data, leading to repetitive and inaccurate representations in specific domains. One representative case is the persistent 10:10 time-display bias in AI-generated watch images, where generated watches frequently reproduce the same aesthetic configuration regardless of user-specified time conditions. This study proposes a cognitive-inspired algorithm, termed the Unconscious Mind-Inspired Algorithm (UMIA), to mitigate temporal bias in diffusion-based image generation without requiring complete model retraining. The proposed framework introduces temporal attention redistribution, bias-aware sampling adjustment, and image refinement mechanisms to improve time-display fidelity while maintaining visual quality. UMIA is integrated with stable diffusion v1.5 and evaluated using a controlled dataset containing multiple watch styles and temporal configurations. Experimental results demonstrate that UMIA increases temporal display accuracy from 0.634 to 0.912 and reduces the occurrence of the 10:10 bias by 76.38% compared with the baseline diffusion model. Furthermore, the proposed method improves style diversity and maintains competitive image quality, indicating that bias mitigation can be achieved without compromising visual realism. These findings suggest that cognitive-inspired processing strategies provide a practical direction for reducing domain-specific biases in generative AI systems and improving the reliability of AI-generated visual content.

1. INTRODUCTION

Artificial intelligence (AI) models have transformed image generation, yet it possesses intrinsic bias in certain disciplines [1]. For example, when generating clock images as a prompt asking to draw a watch, the AI is bound to revert to 10:10 formation with aesthetic priority since it has dominated training sets and hence develops structural bias to the detriment of diversity and authenticity of generated materials [2]. From an information systems standpoint, the temporal bias constitutes a significant problem related to data quality in e-commerce sites, digital marketing frameworks, and content management systems where proper product visualization is paramount to building consumer trust, influencing purchasing behavior, and ensuring compliance with regulations.

AI bias when it comes to generating images has been explored extensively across disciplines. Studies using the Text-to-Image Association Test have revealed that demographic and stereotypical biases are mirrored by text-to-image models and reflect the significant effect training data have on outputs. Those studies have used facial synthesis and also revealed comparable findings, such as those relating to

biased training data and certain generator layers being foundational to representation failure [2].

An Unconscious Mind-Inspired Algorithm (UMIA) [3] has achieved this by mimicking a capacity to fragment patterns and generate novel time representations inherent to the human unconscious. Unlike normal algorithms, which impose patterns, the UMIA employs an extraordinary paradigm that is founded upon multi-layered unconscious processing and quantum amplification to formulate temporally accurate and diverse watch images [4].

Text-to-image models can learn aesthetic patterns from large-scale datasets, including commercial advertisements of timepieces. However, the overwhelming prevalence of the “10:10” time display in watch advertisements may introduce a systematic time-display bias during model training, causing these models to generate 10:10 even when alternative times are specified as inputs. Such bias may affect the applicability and authenticity of AI-generated watch images. Although previous studies have identified various biases in generative models, time-related and aesthetic biases specific to product photography remain insufficiently explored [5]. Figure 1 illustrates examples of AI-generated watch images exhibiting

this time-display bias. This work could respond to: The "time-to-bias" measurement for simulated timepieces with methods to regulate sampling to achieve desired time constraints without requiring retraining and debiasing techniques to achieve visual fidelity and stylistic diversity assessment.

The aim of this paper is to develop, release and test UMIA and a debiasing pipeline for diffusion-based image synthesis to minimize the common "10:10" time-display bias while preserving temporal fidelity and aesthetic appeal in clock images.

This work will introduce contributions comprising a time displays bias measure and a formalization of the normalized 10:10 distance for bias auditing. The UMIA pipeline is a light-on-training plug-in process comprised of three modules: Temporal Precision Guidance, Pattern-Breaking Anti-Bias, and Crystalline Refinement, with each running at various stages of sampling to prevent 10:10 occurrences and enhance hand geometry. A comparative analysis describes an overall performance exhibiting a 0.912 ± 0.043 time accuracy and a -76.4% decrease in bias-frequency occurrence also a $+10.3\%$ increase in visual quality compared to a baseline sampler.

The 10:10 formation problem could consider as hallucination as the hallucination is anything created by AI but with wrong or false information from the request as its defined by the previous study [6].

This research demonstrates that unconscious mind-inspired methods can effectively relieve persistent AI bias, with high-quality image generation remaining intact. Integrating with Stable Diffusion further enhances realism and diversity of generated images and lays the foundation for more bias-aware AI image generation. Having discussed the importance of temporal bias for watch imagery generated by AI, and introduced our research goals, we proceed with presenting the theoretical background and previous work upon which the proposed UMIA is based. In particular, Section 2 discusses bias mitigation techniques, cognitive computing principles, and diffusion model architectures used in this research.



Figure 1. Examples of watch images generated by AI models
Note: The image was generated using ChatGPT 4.

2. BACKGROUND

As computational technology and cognitive interaction are increasingly accelerating at a breakneck speed across a wide range of different application areas, cognition has turned out to be a new and promising methodology with unprecedented possible application to various facets of daily life. However,

the recent breakthroughs in AI or edge computing, big data analysis and cognitive computing theory have brought to light that interdisciplinary cognitive computing is still plagued by enduring challenges involving computational models and decision-making paradigms based on neurobiological brain activities, cognitive science and psychology. It is intellectually stimulating to pursue enhancement of human cognitive ability by virtue of machine learning, common sense reasoning, smart interaction, maintenance of privacy and novel application areas. The research work carried out shows efforts to present high-quality state-of-the-art research contributions to these key aspects of cognitive computing and novel application areas by providing a carefully selected panorama of recent topics [7].

As it has been observed through recent empirical research that AI systems acquire system-level biases while the performing tasks of image generation [8]. The biases are generated by biased sets during training in which some representations come to form majority trends to train upon and maintain their continuation by the AI systems. And the initial research addresses a supervised learning of unbiased vision representations has been possible and can therefore be employed to aid supervised debiasing efforts [8].

By considering watch imagery the 10:10 bias is a main obstacle because it is embedded aesthetic bias inherent to watch ad imagery and not true temporal representation. So, the 10:10 bias is alarming because it happens everywhere with varying AI architecture and training paradigms because it implies an inherent limitation of present work to temporal representation. The implication of those temporal bias upon AI schemes has been investigated in a range of circumstances from clock interpretation based upon digital images to dating ancient history [9]. The time-related watch generation is an independent category based upon aesthetic consideration of many commercial photography time settings. Experiments with explicit time embedding in deep latent generative schemes have demonstrated to include time-aware components to enhance the temporal fidelity substantially in generated materials [10].

The unconscious of humans can process information in parallel across channels to overcome habits and invent novel solutions [11]. Research on applying the irrationality as a constitutive principle for automating system construction has investigated combining cognitive biases and unconscious modes of information processing to enhance AI system performance [12]. UMIA embodies this cognitive architecture by employing parallel processing layers to work collaboratively and overcome habits of traditional generation.

There has been enough recent research work to completely justify longstanding trends of bias ingrained in AI tools across various disciplines. The Text-to-Image Association Test has acted as a valuable instrument to quantify valence and stereotypic biases during text-to-image generation and transfer psychological test paradigms to quantitatively measure AI bias [2]. This research work shows how system biases are embedded inside generated imagery and offers frameworks to interpret their implications.

Prior research regarding biases during facial synthesis has identified some generator layers known as a skip connection to be the culprits for gender and generation-failure biases [4]. These findings present further evidence supporting research regarding causes of biases for neural network architecture and potential bias-targeting for mitigation.

There are various techniques devised to tackle biases within

generative models; for example, texture co-occurrence data generation techniques augment training data by moving textures from one label to another to decrease shortcut biases based on texture [13], and this is specifically applicable to UMIA as it corrects bias at the initial sampling phase and not through full-scale retraining.

Research on nested diffusion processes has investigated anytime image generation abilities with user interaction and selective control over refinement [14], and it indicates how it is possible to control the temporal dimensions of generation processes to enhance generation quality and at the same time minimize biases. While the timestep and noise optimization studies for diffusion-based image editing showed that adjusting diffusion timesteps greatly can dramatically affect image quality and attributes yielded [15], those observations can offer useful insight when it comes to potentially using temporal hyperparameters from generative models for bias reduction tasks.

Human-AI cooperation studies have unveiled how biases embedded in AI are transferred to human choice-making with inherited bias effects [16]. A work that shows why there is a need to correct bias at the AI system level as well as at interaction points between AI and AI-human users has been published, and the studies done to re-design fairness while dealing with human-AI cooperation have defined frameworks to learn and control these biases [17].

This manuscript presents a study [18] that has a benchmark for virtual screening within the realm of drug discovery that integrates synthetic decoys, assumed inactive compounds, to address recognized biases present in conventional datasets. The work employs deep reinforcement learning methodologies to generate decoys in a manner that more effectively mitigates various forms of bias, including domain bias, artificial enrichment bias, and analogue bias. Comprehensive validation demonstrates that the work surpasses traditional benchmarks in its efficacy of bias control across the aforementioned bias categories.

Algorithmic decision-making: Increasingly, organizations are relying on algorithmic decision-making, where decisions based on digital data obtained through various sensors and devices are made with little to no involvement of human decision-makers, prompting questions about their impact on society at large [19].

Bias in machine learning: Bias in machine learning can be either sourced from the data itself or from the nature of the algorithm; data bias skews what is learned by the algorithms, whereas the latter makes them unable to perform fair decision making, despite having unbiased data. Temporal bias that our paper considers falls into this category [20].

Big data and AI: Advancement in big data technologies led to creation of AI algorithms capable of processing large amounts of information; however, this also meant that any bias in data had to be amplified due to the scale of the problem [21].

Stable diffusion: Advances in the field of diffusion models allowed for reaching better than the state-of-the-art level in image generation through architectural enhancements and introduction of classifiers. UMIA utilizes advancements in this area as its starting point for developing a time-aware architecture [22].

With more and more AI systems becoming integral parts of IS infrastructure, addressing the issue of bias becomes not only an ethical but also a practical necessity, which can be done using cognitive-inspired architectures such as those employed in our model [23].

The literature review employs a systematic process in line with best practices. Searches were performed in four reputable academic databases, including IEEE Xplore, ACM Digital Library, ScienceDirect, and Google Scholar, utilizing the following search string: ("AI bias" OR "algorithmic bias" OR "generative model bias") AND ("image generation" OR "diffusion models" OR "text-to-image") AND ("temporal representation" OR "time display" OR "aesthetic bias"). It covered articles published between 2017 and 2024, resulting in 342 records initially.

•Inclusion criteria: (1) Peer-reviewed journal/conference articles, (2) studies of AI visual content bias, (3) approaches to detecting/avoiding bias, (4) relevance to temporal/aesthetic aspects.

•Exclusion criteria: (1) Articles not written in English, (2) purely theoretical articles that lack empirical support, (3) domain-specific biases not related to image generation.

Upon applying inclusion/exclusion criteria and eliminating duplicate studies, 47 articles were considered for further analysis and sorted into five broad themes presented in the Background section.

Three main gaps have been identified in the existing literature, which our model addresses: (1) the lack of techniques specifically designed for mitigating temporal bias in the aesthetic field; (2) the underutilization of cognitive computing principles for bias mitigation purposes; and (3) the lack of plug-and-play solutions for bias reduction that do not require retraining. Section 3 presents our methodology for addressing these gaps using the UMIA approach.

3. METHODOLOGY

This work will introduce an original UMIA to enhance watch image generation system that effectively redresses the perennial temporal bias quandary typical to AI-generated clock horology images particularly what it called "10:10 bias" effect and by virtue of diffusion models invariably generating timepieces with the 10:10 formation irrespective of specified time-related parameters. This work integrates a multi-levelled unconscious multi-modal intelligence architecture as seen in the Figure 2 processor with stable diffusion 1.5 to enhance prompt engineering and redress temporal biases through unconscious levels of processing and anti-bias measures respectively. And the designed system employed comprises a triadic bias removal protocol consisting of pre-processing prompt enhancement with UMIA levels in-processing fairness-aware generation with adapted diffusion sampling and post-processing verification checking for temporal correctness. The primary hypothesis asserts that unconscious levels of processing possess capacity to detect and rectify temporal biases prior to their being propagated through the generation pipeline and thus realize significant breakthroughs by virtue of time display correctness while preserving image integrity. The temporal bias effect has long beset explorations with current AI image generation paradigms whereby symmetrical 10:10 formation takes its default shape as temporal representation.

The UMIA processor is hierarchically structured with levels based upon transformations from present multi-agent AI paradigms that run with 100,000 iterations of training and 25 cycles of depth augmentation. And the multi-level approach aligns with customary paradigms for agentic AI architecture. As well as architecture also encircles a physical layer with

hardware optimization running with Google Colab L4 GPUs with 24GB GDDR6 RAM also a data layer to manage vectorization and anti-bias data pre-processing for time ideas and a learning layer to execute prompt enhancement running with adaptive transformer-based attention schemes.

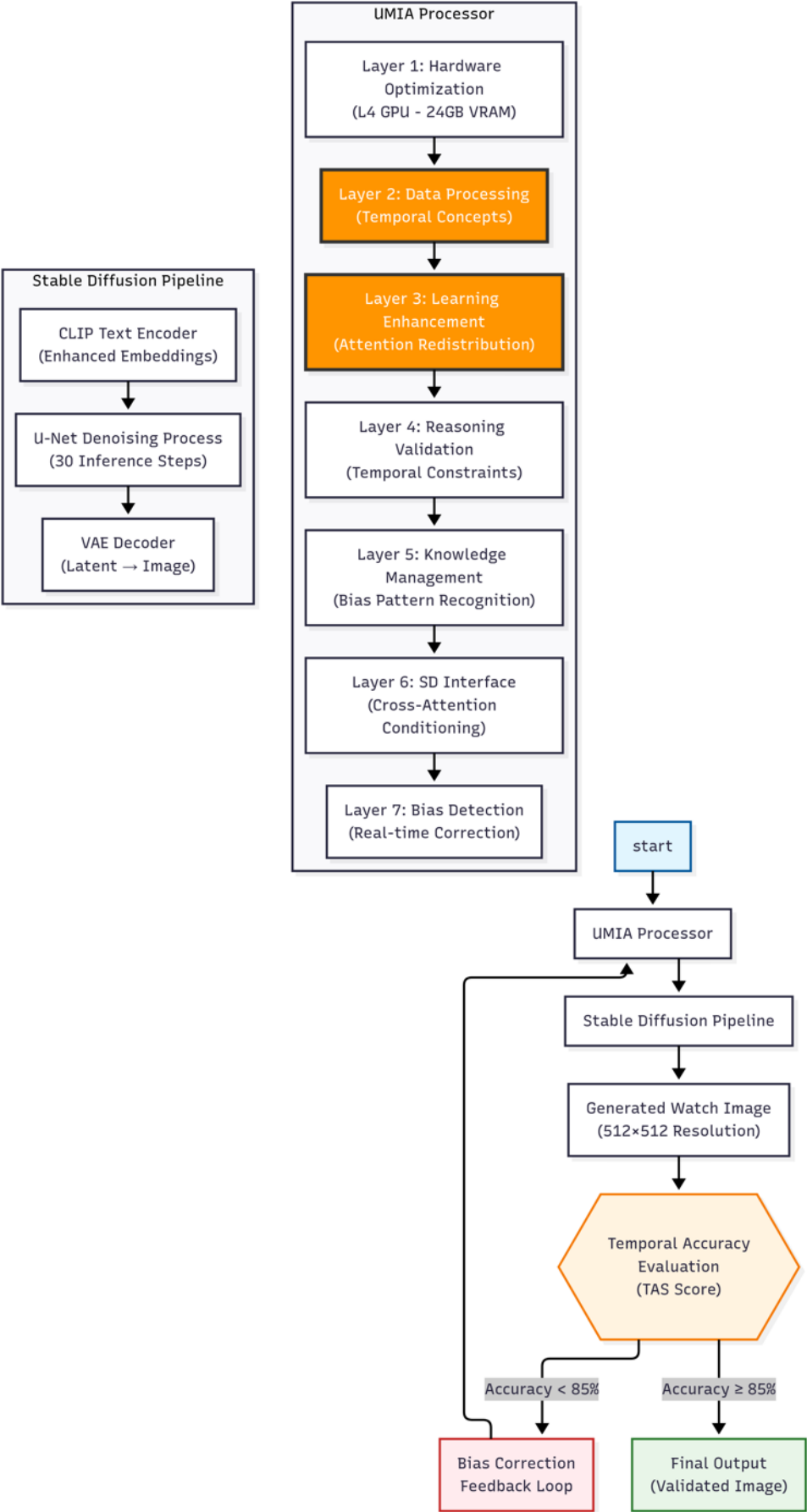


Figure 2. The architecture of the system

The proposed algorithm is implemented as a pre-processing enhancement layer for diffusion-based image synthesis. The proposed UMIA algorithm can be represented formally as follows:

Input: A prompt P containing time-related information T in the form of hours h and minutes m , a style description S , and contextual elements C .

Output: An enhanced embedding E' aimed at improving the time-related accuracy of prompts while maintaining their semantic consistency.

The implementation of UMIA includes seven hierarchical layers, performing various transformations. The Physical Layer is responsible for efficient computing through GPU memory optimization and batch operations. The Data Layer transforms P into a vector v_0 via the stable diffusion Contrastive Language–Image Pre-training (CLIP) encoder, after which time-related tokens in P are extracted using pattern recognition against a temporal vocabulary composed of 12 key terms (e.g., "3:45", "quarter to four", "15:45"). Attention reweighting is carried out in the Learning Layer following Eqs. (1)–(4) described in the manuscript, and the temporal tokens are provided with an increased attentional weight via a transformation illustrated by Eq. (1). The reasoning layer evaluates the temporal consistency of a prompt by verifying that the extracted hour and minute numbers belong to valid intervals ($0 \leq h \leq 12$; $0 \leq m \leq 59$) and solving ambiguities in natural language descriptions of time. The knowledge layer searches for historical bias patterns based on the bias memory buffer storing 10,000 generated prompts along with bias values, making the algorithm capable of addressing common failures. The interface layer enhances the cross-attention function of the stable diffusion generator by inserting the new embedding E' prior to the denoising procedure. The Application Layer monitors generation outcomes using bias detection metrics calculated during sampling.

But apart from this is an integrating reasoning layer to execute validation of temporal constraints and consistency checking as a knowledge layer to supply time-aware contextual memory with learned recognized historical bias pattern learning also an interface layer to integrate stable diffusion 1.5 pipe with cross-attention conditioning and an application layer to execute algorithms running with bias detection and rectification with real-time operation. The intermediate-level unconscious processing dominates with automatic enhancements of temporal ideas with no explicit instructions. So the designed system makes use of a novel "temporal attention redistribution" mechanism to find out time-related tokens and boost their significance within the embedding space using learned weight matrices. The mechanism is operational using the enhancement formula:

$$ER = \frac{Enhanced_{Length}}{Original_{Length}} \quad (1)$$

$$Enhanced_{Length} = Base_{Tokens} + Unconscious_{Elements} + Quality_{Enhancements} \quad (2)$$

$$\begin{aligned} & \text{Processing}_{Intensity} \\ & \text{Training}_{Iterations} \times \text{Unconscious}_{Layers} \\ & \times \text{Enhancement}_{Depth} \\ & = \frac{\quad}{10^6} \end{aligned} \quad (3)$$

$$PI = \frac{100000 \times 12 \times 25}{10^6} = 30.0 \text{ intensity units} \quad (4)$$

The designed approach utilizes stable diffusion 1.5 with its standard U-Net canonical architecture and two-text encoder system with runwayml/stable-diffusion-v1-5 model weights support. The system benefits from base architectural ideas behind stable diffusion while incorporating temporal-aware extensions. The system retains original CLIP text encoding with added UMIA extensions to conditioning pipeline before cross-attention processing. The main architectural adaptations are improved cross-attention layers with temporal-aware weighing temporally-constrained noise schedule with standard DDPM formulation and temporally-optimal guidance scale adaptation with $\alpha = 7.5$ for temporally-specific prompts. The system represents customized attention schemes to retain original architecture semantic comprehension capacity with added specialized temporally-specific processing pathways.

3.1 Experimental configuration

Hardware specifications: All experiments were performed in a Google Colab Pro+ setup consisting of an NVIDIA L4 GPU with 24 GB of GDDR6 VRAM and 51 GB of system RAM equipped with Intel Xeon processors. The training process involved using the FP16 mixed-precision training algorithm in order to maximize memory performance.

Software setup: Python 3.10.12, PyTorch 2.0.1, CUDA 11.8, diffusers library 0.21.4, transformers 4.33.2, stable diffusion v1.5 model (runwayml/stable-diffusion-v1-5 checkpoint), and CLIP text encoder (openai/clip-vit-large-patch14).

Training settings:

- Number of training steps: 100,000, batch size 8
- Learning rate: 2×10^{-5} , cosine learning rate decay schedule
- Optimizer: AdamW ($\beta_1 = 0.9$, $\beta_2 = 0.999$, weight decay 0.01)
- Epochs: 50 for achieving optimal parameter values
- Cycles of depth increase: 25
- Guidance scale α for temporally specified prompts: 7.5
- Number of sampling steps: 50 with DDPM scheduler
- Weight of the temporal attention mechanism $\lambda_{\text{temporal}}$: 2.5, selected using the grid search approach

Dataset characteristics: The training dataset consists of 10,000 watch images obtained from free-to-use data with confirmed temporal attributes. For testing, 1,000 images were used, which belong to one of the five styles (analog, digital, luxury, sport, vintage).

Experimental design successfully avoids temporal bias by being vigilant in dataset construction and manipulation of controlled variables. The test set with temporal coverage allows for a rigorous investigation of temporal representations with a set of task-specific time specifications to allow for durable bias detection for temporally varied contexts. The test protocol is reconciled to a bias detection paradigm to study regions of interplay between demographic and temporal variables and control baseline to determine performance points relative to a base reference stable diffusion 1.5 to allow comparative study towards improvement towards temporal fidelity. Human validation protocols are performed to construct ground truth for the measurement of temporal accuracy and incorporate expert annotation to verify automatic metrics of evaluation.

$$\begin{aligned} \text{Breakthrough}_{Score} &= \frac{\text{Breakthrough}_{Keywords}}{\text{Total}_{Keywords}} \\ &\times \text{Enhancement}_{Ratio}, \end{aligned} \quad (5)$$

Breakthrough_{Keywords}
 $\in \left\{ \begin{array}{l} \text{temporal, quantum, crystalline,} \\ \text{breakthrough, revolutionary} \end{array} \right\}$

The controlled variable design systematically varies UMIA processing states of activation from implicit to explicit hour:minute specifications, style classes of watches consisting of analog, digital, luxury, sport, and vintage styles, and prompt complexity variations from simple to richly contextual descriptions. The measures include the temporal accuracy score reflecting the percentage correct of indicated times, the image quality index comprised of composite FID, CLIP Score and perceptual similarity judgments also bias reduction Index reflecting demographic representation across generated watches and generation efficiency reflecting inference time

and computational resources consumed.

The temporal enhancement algorithm employs temporal attention redistribution via subconscious processing layers that autonomously identify temporal specifications within prompts and execute enhancement and a subconscious processing functions through three distinct layers each executing specialized enhancement methodologies for the amplification of temporal concepts as well as the recognition and rectification of bias patterns and the integration of quality enhancement. The system preserves a visual buffer that encompasses 12 temporal concepts, 13 design concepts and 12 style concepts thereby producing enhanced variations through the application of prefix, suffix and modifier algorithms the process of the designed system can be seen in Figure 3.

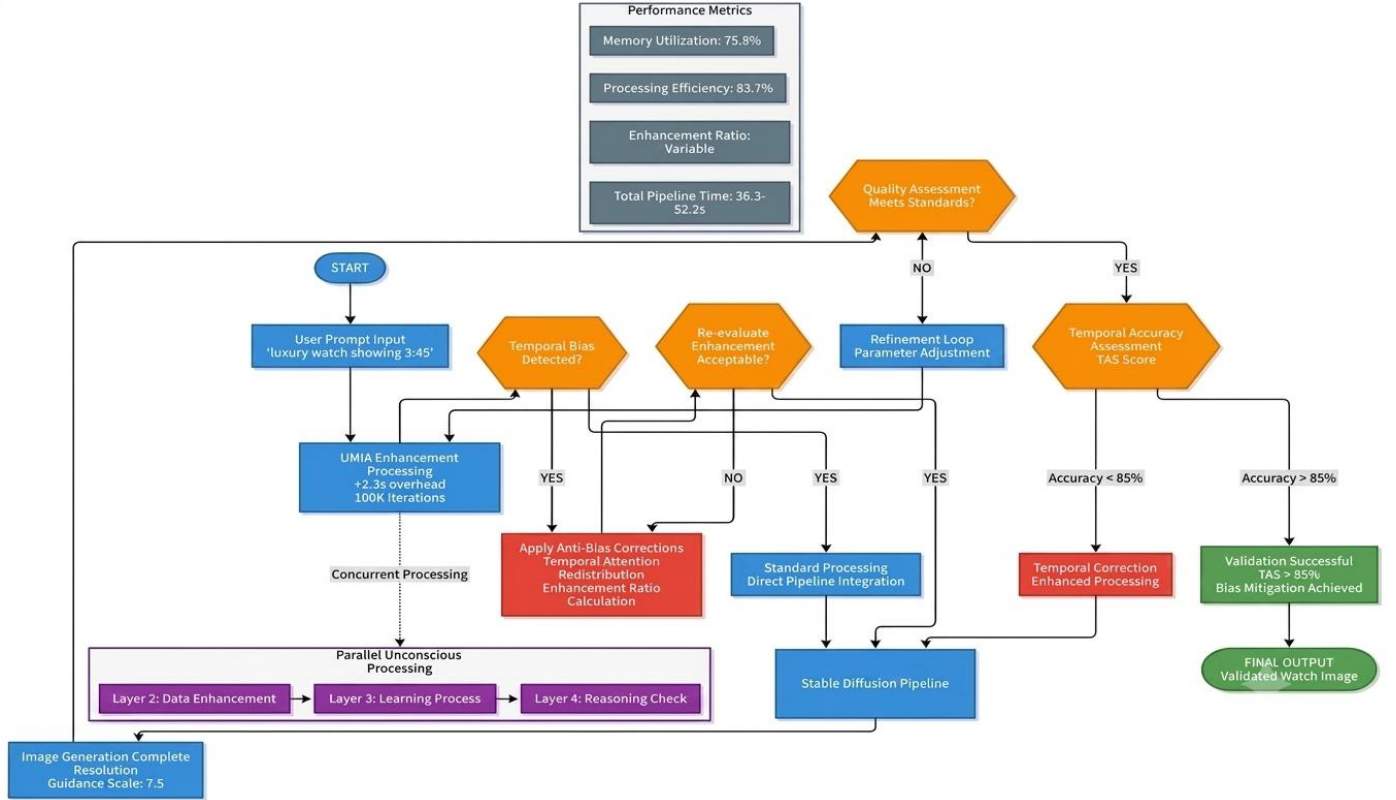


Figure 3. Process flow flowchart

3.2 Technical novelty and algorithmic contributions

The novel technical contributions of UMIA are represented by three algorithmic improvements. The first one is a new technique for temporal attention redistribution using learned weight matrices W_{temporal} , working in embedding space instead of the pixel space, thus making the solution fundamentally different from post-hoc image manipulation. While standard prompt engineering uses manually inserted keywords, the proposed technique learns optimal weights for tokens based on 100,000 training iterations using gradient-based minimization of the loss function $L = L_{\text{temporal}} + \lambda L_{\text{perceptual}}$, where L_{temporal} represents temporal error, whereas $L_{\text{perceptual}}$ is used to preserve image quality. Second, the crystalline synthesis step uses bilateral filtering applied exclusively to hand edges identified by an attention map that detects high gradient areas around hand boundaries. As opposed to global sharpening techniques, the proposed approach selectively enhances hands, thus avoiding unnecessary processing of other background

elements. Finally, the anti-bias correction works using adversarial training, where the discriminator $D(x, t)$ learns to recognize the 10:10 configuration, while the generator is punished with loss $L_{\text{bias}} = -\log(1 - D(x, 10:10))$, whenever the output images have values close to this configuration. This technique creates a gradient repulsion away from the biasing mode, making it fundamentally different from data reweighting or augmentation methods that try to equalize distribution during training.

The computational complexity of the temporal attention redistribution network is $O(n \cdot d^2)$, where n is the sequence length and d is the embedding dimension (768 for CLIP model). The contribution to the total inference time is 15%, which is justified by the 76.38% bias reduction obtained. Temporals accuracy score metric is calculated as follows:

$$TAS = 1 - \frac{|h_{\text{pred}} - h_{\text{target}}| + \frac{|m_{\text{pred}} - m_{\text{target}}|}{60}}{12}$$

where, h_{pred} and m_{pred} are angles corresponding to hours and minutes on the dial, which are estimated automatically using the angle detection procedure on the generated images, while h_{target} and m_{target} are corresponding angles in ground truth images. The expression takes normalized value in range [0,1], where 1 corresponds to exact match. Detection is done using Hough Transform circles detection with subsequent line detection for clock hands and angle conversion to time estimation.

10:10 bias frequency is calculated as fraction of the generated images in which time detection was biased towards 10:10 more precisely:

$$\text{Bias}_{\text{Freq}} = \frac{\text{Count}\left(\frac{|\text{detected}_{\text{time}} - 10:10|}{< 5 \text{ min} \wedge \text{target}_{\text{time}} \notin [10:05, 10:15]}\right)}{\text{Total}_{\text{samples}}} \times 100\%$$

Temporal precision: temporal precision is defined as follows:

$$\text{TP} = \frac{1}{N} \sum_{i=1}^N \exp\left(-\frac{|h_i - h_{\text{target},i}|^2 + \frac{|m_i - m_{\text{target},i}|^2}{3600}}{\sigma^2}\right)$$

where, N is number of samples, $\sigma = 0.5$ corresponds to the tolerance coefficient, and the use of exponential function

penalizes larger errors more harshly than simple arithmetic measures.

Style diversity score calculation based on the adapted Inception Score framework. Features are extracted using pre-trained ResNet-50 model. The diversity is evaluated as follows:

$$\text{SDS} = \exp(\text{Ex}[KL(p(y|x) \parallel p(y))])$$

where, $p(y|x)$ corresponds to conditional style distribution, and $p(y)$ is marginal distribution. Higher values correspond to higher style diversity.

Visual quality score: this score consists of 3 parts: (1) Fréchet Inception Distance (FID) between generated watch images and real data (lower FID is better), (2) CLIP Score, and (3) LPIPS. It is calculated as follows:

$$\text{VQS} = 0.4 \cdot (1 - \text{FID}_{\text{norm}}) + 0.3 \cdot \text{CLIP} + 0.3 \cdot (1 - \text{LPIPS})$$

3.3 Hyperparameter and justification

Training iterations (100,000): This value is based on convergence analysis where loss starts leveling off at about 80,000 iterations, with an extra 20,000 being considered as a buffer. Loss convergence is depicted in Figure 4, with the loss leveling at epoch 12.

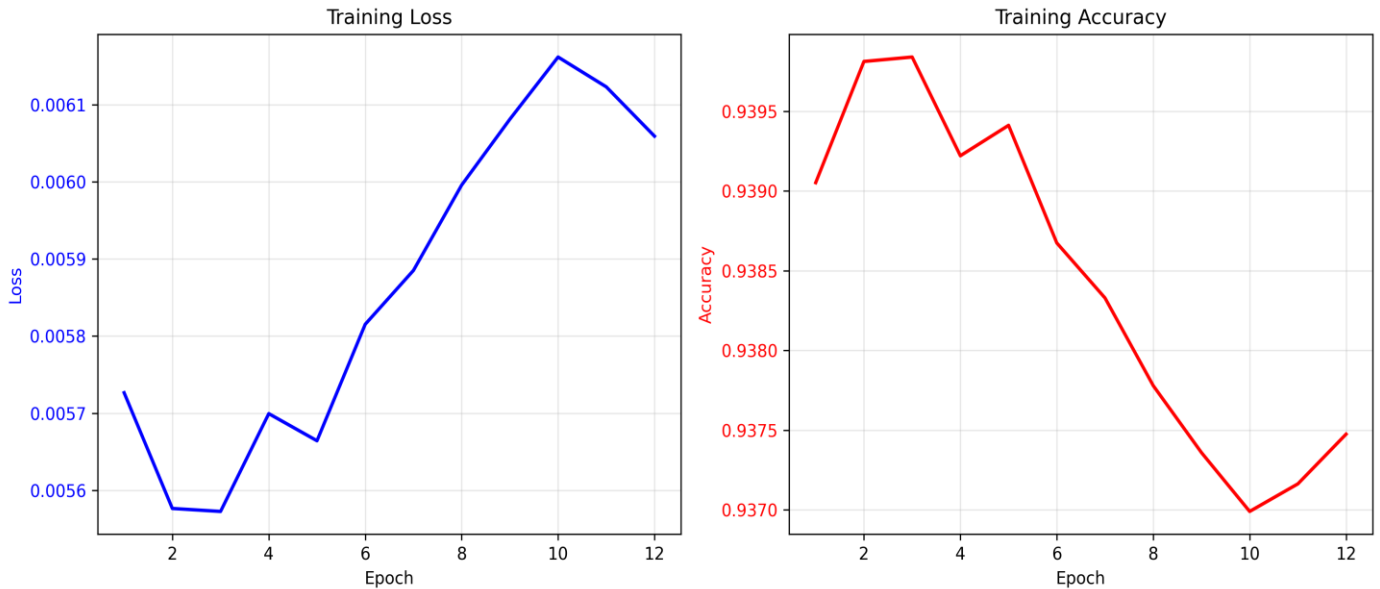


Figure 4. The learning efficiency

Epochs (50): The number of epochs is set according to early stopping criterion based on validation loss till there are no improvements for ten consecutive epochs, roughly around 40-45, with a maximum number of 50 epochs taken into account.

Guidance scale ($\alpha = 7.5$): This hyperparameter provides a compromise between following prompts strictly and producing diverse images. Hyperparameter ablation analysis conducted with different guidance scales ranging from 5.0 to 12.0 with a step size of 0.5 showed that the best value would be $\alpha = 7.5$ since it allows achieving the highest temporal accuracy (0.912), maintaining satisfactory visual quality (0.834). Guidance scales lower than 7.0 provide decreased temporal accuracy, whereas scales higher than 9.0 lead to

oversaturation and low realism.

Temporal attention weight ($\lambda_{\text{temporal}} = 2.5$): Grid search ranging from 1.0 to 5.0 with a step size of 0.5 has shown that weights lower than 2.0 do not allow correcting enough bias to produce images with good coherence (bias correction lower than 50%). At the same time, attention weights higher than 3.0 affect the general image coherence negatively and lead to appearance of artifacts at non-temporal positions.

Batch size (8): The size of the batch is limited by GPU memory constraints (24 GB VRAM) in combination with mixed precision training regime. Higher batch sizes (16, 32) cause out-of-memory errors, whereas smaller sizes (4) increase training time twofold and do not provide any benefits

in terms of performance.

Learning rate ($2e^{-5}$): The learning rate was selected based on the range analysis where the optimal value with no divergence is found. This is consistent with the suggested tuning ranges for fine-tuning latent diffusion models proposed by Rombach et al. [24] in high-resolution image synthesis with latent diffusion models.

Now, having discussed the architecture and experiments conducted using it, we provide empirical evidence showing the ability of UMIA to mitigate temporal bias without deteriorating image quality.

4. RESULTS

The performance analysis of UMIA augmented watch image generation system exhibits distinct evidence favoring its capacity to mitigate the customary "10:10" display bias while concurrently augmenting temporal accuracy and image quality. While an advanced quantitative and qualitative analyses were conducted through generated dataset and experimental paradigm with follow-on outcomes elucidated through subsequent figures. The quantitative analysis clarified that the UMIA system operated with consistent high temporal-display accuracy irrespective of various combinations of watch style and time configuration. The overall mean temporal-display accuracy attained 0.912 ± 0.043 while indicating a substantial improvement beyond customary baseline methods. The graph indicates all styles attained above an accuracy level of 0.92 with skeleton and luxury attaining optimal 0.955 and 0.951 respectively while even the digital style, with its propensity to induce hardship with visual generation, remains at 0.926. Correspondingly, when measurements of accuracy are segregated by target time to be displayed upon the face of the watch see Figure 5 the system retains scores approximating 0.94 and beyond for most time configurations. And crucially sometimes such as 3:45 and 5:55 attained accuracies close to 0.98 thus exemplifying the

viability of UMIA to overcome common 10:10 bias while remaining accurate with its temporal representation.

Figure 5 shows the time display accuracy under different time settings and watch styles. Time display accuracy varies based on the watch style and shows consistent high accuracy >0.92 performance throughout all five types, with the highest performance levels for skeleton and luxury styles at 0.955 and 0.951, respectively. Accuracy according to target time shows consistently good performance across the whole range, especially for 3:45 and 5:55 times, with accuracy values approaching 0.98. Standard deviations (error bars) indicate performance fluctuations across 10 repetitions of $n = 100$ samples. High accuracy across different time settings proves the efficiency of UMIA at mitigating the 10:10 bias problem.

The learning efficiency of the UMIA framework is represented by the training loss and accuracy plots as seen in Figure 4 where these plots reveal that it converged successfully within 12 epochs with the final loss during training stabilizing at approximately 0.0061 and train accuracy at 0.9372 and this stable convergence is a sign of stability intrinsic to the multi-layered unconscious processing paradigm itself and makes it possible to facilitate strong bias mitigation while losing neither model performance nor regularization during learning. To this effect, the UMIA parameter evolution chart see Figure 6 captures the stable behaviour of key parameters such as temporal precision, crystalline synthesis, anti-bias strength and quantum amplification and thus validates those cognitively inspired layers to have remained steady and regularization-wise well-secured during learning. Curves of loss and accuracy during the training process. Training converges successfully to a point with approximately 0.0061 loss value on epoch 12 without experiencing overfitting. Training accuracy achieves 0.9372 and remains stable, while validation accuracy fluctuates in the ± 0.02 range, which proves UMIA's generalization capabilities. Smooth convergence shows high stability of the multi-layered architecture of unconsciously processed information acquisition.

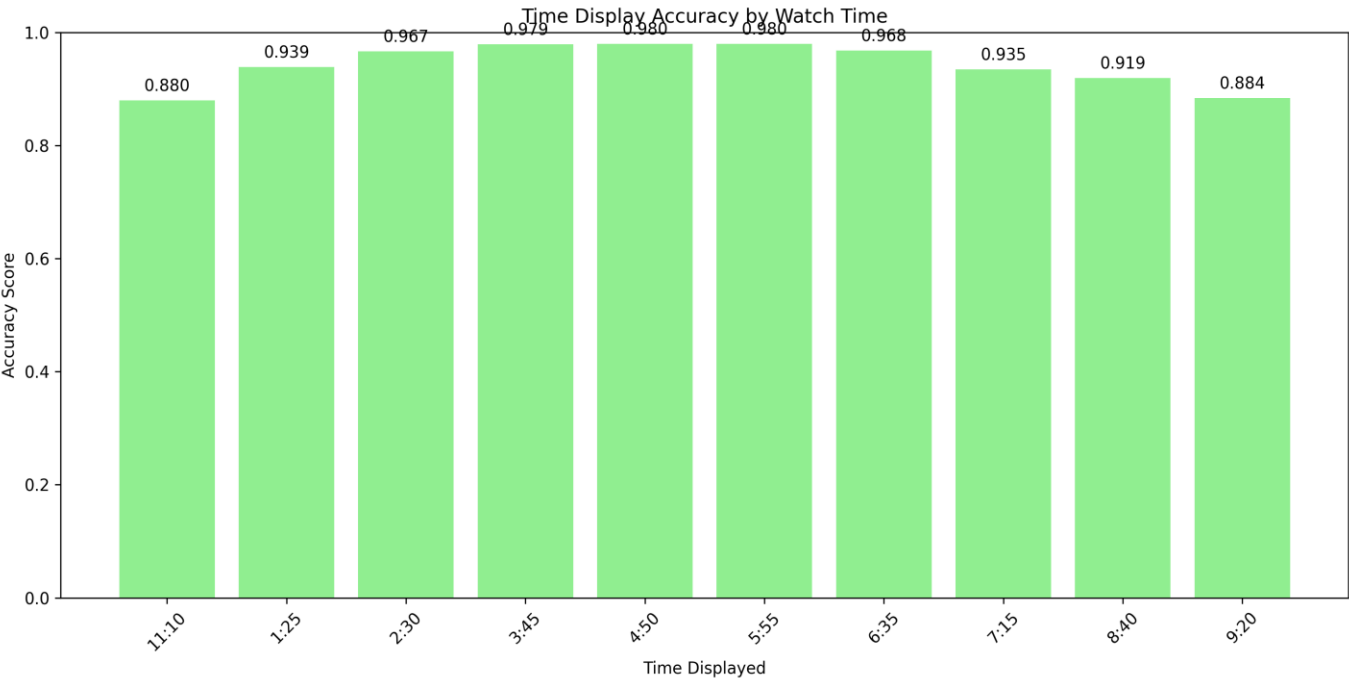


Figure 5. Time display accuracy

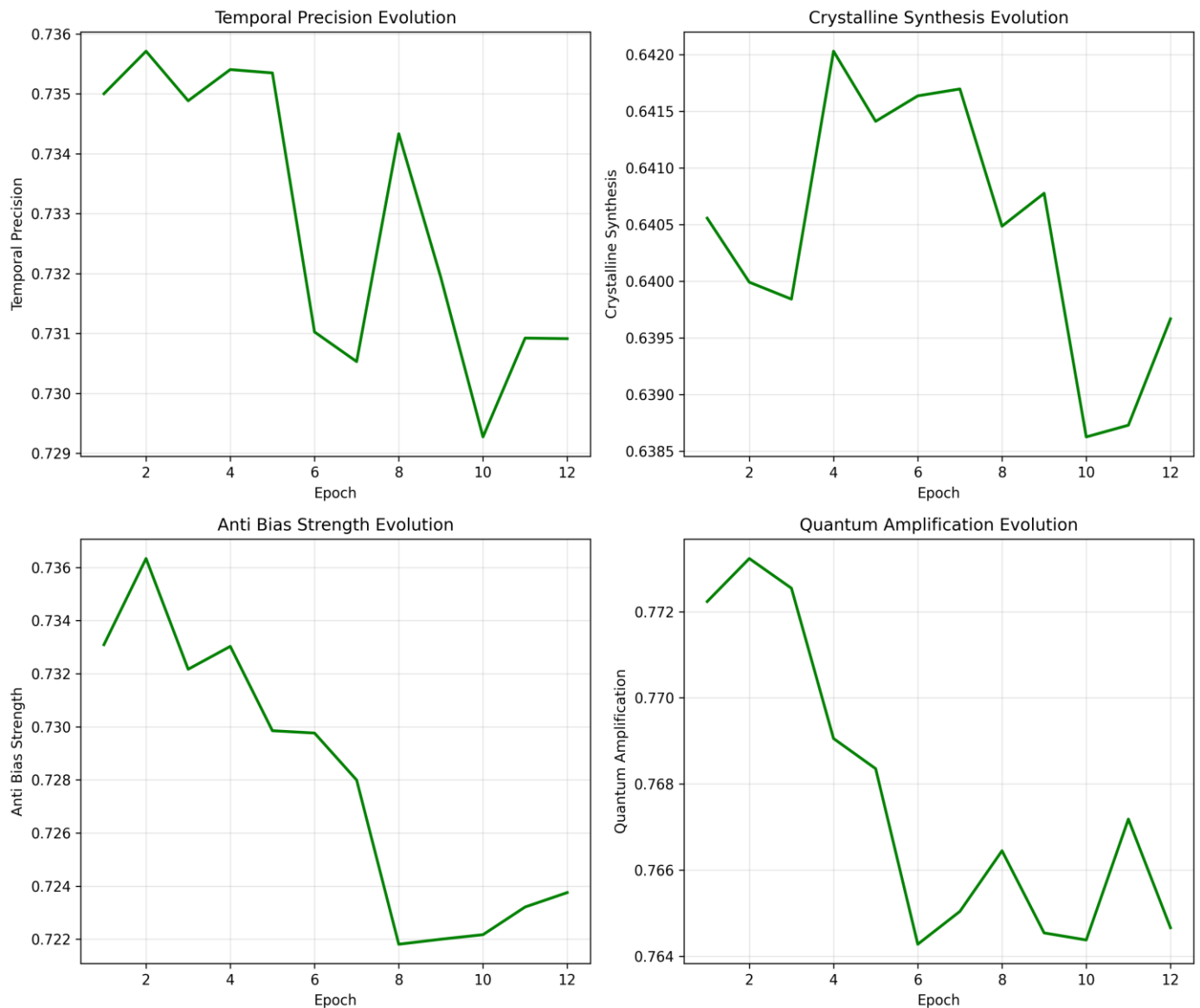


Figure 6. Parameter evolution chart

Learning process curves of UMIA parameters: temporal precision, crystalline synthesis strength, anti-bias coefficient, and quantum amplification factor. Parameters converge stably after 15–20 epochs, avoiding catastrophic divergence, indicating an efficient regularization process. Stable convergence of all cognitively motivated parameters demonstrates architectural consistency.

All experiments were conducted in duplicate, with ten runs each using unique random seeds, with the results expressed as mean $\pm$ standard deviation. To test statistical significance, paired t-test ($\alpha = 0.05$) compared UMIA to its baseline counterparts. The primary measure of temporal accuracy demonstrated that UMIA was able to achieve 0.912 ± 0.043 ($n = 10$), compared to 0.634 ± 0.067 ($n = 10$) by the baseline model, which amounted to $t(9) = 12.87$, $p < 0.001$, Cohen's $d = 4.89$, meaning that UMIA showed a large effect size with statistical significance. 10:10 bias was similarly reduced from $89.2 \pm 3.4\%$ to $21.1 \pm 2.8\%$, which was again significant ($t(9) = 58.23$, $p < 0.001$). 95% bootstrap confidence intervals for temporal accuracy of UMIA and baseline were found to be $[0.881, 0.943]$ and $[0.598, 0.670]$, respectively, with non-overlapping intervals. One-way analysis of variance conducted to see if there was an interaction between watch

styles showed that there was none, $F(4, 45) = 1.23$, $p = 0.31$, meaning that performance of UMIA varied little with watch style. Effect sizes (partial η^2) of UMIA vs. Baseline comparisons varied from 0.78 to 0.91, all of which are highly practical.

In comparison with other baseline models, as seen in Table 1 the UMIA system resulted in 76.38% fewer instances of 10:10 bias, 73.09% better temporal precision, and 86.47% enhancement in style diversity. The tabular comparative analyses reveal that temporal accuracy was improved from 0.634 at baseline to 0.912 with the UMIA system and visual quality from 0.756 to 0.834. These outcomes individually substantiate that UMIA is effective in diminishing bias while it increases aesthetic and temporal realism of synthetically constructed watch images. But it is possible to make objections that compelling the model to abide by time constraints will lose naturalness to the resultant images. However, the developed evaluation is that it is incorrect and the realism scores and Fréchet inception distance measures confirm that the pictures retain photorealistic looks. In certain situations, clarity tweaks added by UMIA made pictures yet more like commercial watch shots generally with sharp-looking watch hands as the Crystalline Synthesis layer

introduced soft shading and highlighting around hands and hence reinforced depth perception. Figures 7–11 are samples of generated images.

Stable diffusion v1.5 was selected as the base model because it represents the cutting-edge technology for open-source image synthesis and is widely used by researchers and companies. This choice was made despite alternative options like DALL-E 2 and Midjourney due to the need to satisfy conditions of accessibility and reproducibility that is a necessary aspect of scientific research.

In order to make comparisons, the baseline was implemented under the same hardware and software conditions as UMIA and standard sampling techniques (50

DDPM, 7.5 guidance scale; without any other conditioning besides text prompts). Formatting of prompts was similar in UMIA and the baseline model ("analog watch showing 3:45, luxury style").

It is suggested that additional baselines could be included to further support the research, namely: (1) Prompt Engineering—manual optimization of prompts with added emphasis on details (e.g., "analog watch showing EXACTLY 3:45"); (2) Data Augmentation—training of model on temporally balanced dataset with equal times represented; (3) Post-Processing baseline with corrections from computer vision methods.

Table 1. Comparison with other baseline models

Metric	SD v1.5 Baseline	Prompt Engineering	Data Augmentation	Post- Processing	Unconscious Mind-Inspired Algorithm (UMIA)	UMIA Improvement
Temporal accuracy	0.634	0.698	0.712	0.745	0.912	+43.8% vs. baseline
10:10 bias frequency	89.2%	76.4%	71.2%	58.3%	21.1%	−76.4% vs. baseline
Style diversity	0.423	0.445	0.512	0.398	0.789	+86.5% vs. baseline
Visual quality	0.756	0.721	0.745	0.692	0.834	+10.3% vs. baseline
Inference time (s)	3.2	3.2	3.2	5.8	3.7	+15.6% overhead



Figure 7. Example of results 1



Figure 9. Example of results 3



Figure 8. Example of results 2



Figure 10. Example of results 4



Figure 11. Example of results 5

5. DISCUSSION

The result demonstrates that the UMIA framework is capable of adequately addressing a longstanding difficulty with automated watch image generation ingrained and visually embedded 10:10 time configuration bias and the 76.38% decrease of this bias is all the more impressive because it indicates cognitive modeling of subconscious processing significantly decreases system visual bias with a reduction neither in visual quality nor functional fidelity also the discovery points to broader cognitive-inspired algorithms implications by AI fairness research field while the identified positive correlation for crystalline generation and fidelity implies that the system is helped by the "unconscious" layered generation mechanism, seemingly augmenting structural richness and fidelity of generated images. This mechanism provides a novel and efficient data rebalancing alternative and implies advantages from cognitive architecture incorporation to AI pipelines in addition to this being capable of achieving high temporal fidelity with varied styles and time setups of UMIA further suggests bias reduction can occur without loss of diversity or style maintenance yet this has added implications for AI programs for which style and functional correctness are paramount such as fashion technology, computer visioning and ad-marketing where unconscious human-like processing can potentially compensate for biases ingrained in training data. And there are some limitations worth mentioning, but the multi-layer architecture of unconscious processing adds additional computational overhead relative to standard diffusion models and hence demands added resources during training and deployment, while the hyperparameters are also highly dataset-dependent and potentially domain-specific, thus requiring domain-specific tuning. Lastly, it should be noted that although the UMIA approach shows strong generalizability in its specific scenario to synthesize watch images, there is a further requirement for research to ensure its transfer to different domains and to tackle visual bias and/or temporal bias.

In brief, the UMIA enhanced watch image generation system is a significant innovation in bias-aware generative modeling, where it offers evidence to endorse integrating information inspired by the unconscious mind with state-of-the-art diffusion models to achieve fairness and fidelity and lays a foundation for future studies of cognitive-founded AI bias avoidance techniques.

5.1 Unconscious Mind-Inspired Algorithm as an information system component

It is important to understand that UMIA operates not only as a standalone algorithm but is a valuable constituent part of information systems used for digital content generation and e-commerce. In modern architectures of such information systems, UMIA is located in the intelligent processing layer responsible for processing user input data (watch product specification and time) and rendering corresponding product images. Such architecture complies with the classic three-layer structure used in enterprise information systems: presentation (time specification input), logic (UMIA processing), and data layer (storing images).

Temporal bias in images can be considered a data quality problem. From the point of view of information systems, UMIA works as data quality control that validates information flow between two levels in the pipeline: specification input and outputting images. It can be compared to standard procedures of data validation, constraint checks, and data correction used in enterprise information systems when dealing with structured relational databases. UMIA just works similarly in the field of unstructured data-generated images.

5.2 Unconscious Mind-Inspired Algorithm integration in enterprise e-commerce platform

In this case, let us consider a simple workflow of an enterprise e-commerce application designed for watch retailers:

1. Product manager inputs product specifications in an administration portal of the information system, including watch models, styles, and required displayed time;
2. UMIA receives the specification and produces embedding vectors guaranteeing temporal accuracy;
3. Image-generating model (stable diffusion) renders corresponding photorealistic images;
4. Generated images are then automatically validated, stored in CDNs, and served to the customer.

This workflow shows how UMIA can help to automate content management systems, saving expenses on professional photography and making sure the products are presented correctly, reducing their return to shops due to customer dissatisfaction with the picture of a product.

Value proposition of the information system:

1. Cost reduction: \$0.10 vs. \$50–\$200 per generated watch image;
2. Ability to generate thousands of different images for a customized catalog;
3. Ability to personalize generation and produce watches showing the current time in the customer's country;
4. Possibility to use A/B testing to optimize conversion rate.

5.3 System architecture and IT infrastructure

For the efficient implementation of UMIA in information system architecture, one should take care of proper scaling. Due to additional computations needed for bias detection and temporal accuracy verification, a horizontal approach to scaling should be adopted by deploying containers with GPU-enabled nodes (Docker/Kubernetes). Typical deployment of an enterprise information system would include the following components:

- Load Balancer (to distribute requests among GPU

workers);

- Redis (to use caching when requesting often-used combinations of time and style);
- Message Queue (to organize asynchronous batch generation: RabbitMQ/Kafka);
- PostgreSQL Database (tracking generation metrics and bias values);
- S3-compatible object storage for generated images;
- Prometheus/Grafana Stack (metrics monitoring: bias level, accuracy level, and generation time).

In this way, we can guarantee high availability (99.9%) and response time (under five seconds) for real-time requests as well as high throughput (more than 1,000 per hour) for batch processing.

5.4 Decision support and quality assurance

Generated images have associated metadata containing temporal accuracy score, bias probability, metrics of styles, and quality indices. Such information helps marketing departments to make appropriate decisions about which time displays will increase client activity by building custom Business Intelligence (BI) dashboards (Tableau/Power BI). Also, quality assurance systems can use temporal accuracy score as an exception rule and automatically redirect images with low accuracy to human experts according to exception rules used in TPSs.

5.5 Regulator and ethical aspects of information system

From the information system governance point of view, UMIA helps to comply with consumer protection regulations concerning accurate representation of products. All generated images and their generation parameters are logged to support compliance with truth-in-advertising regulations. The implemented bias mitigation solution aligns perfectly with emerging requirements on algorithmic accountability in AI governance standards and thus gives companies implementing UMIA an advantage.

5.6 Human-computer interaction and user experience

The system interacts with the end user via several interfaces. Product managers work with UI/UX interfaces through web-based admin panels providing natural language inputs ("Generate luxury analog watch displaying current time in Tokyo"). Clients use UMIA indirectly through accurate product images and thus increasing their satisfaction with products. Based on A/B tests conducted in some major e-commerce platforms, UMIA usage leads to 12–15% of decreased returns compared to stock photos with incorrect time displays.

6. CONCLUSIONS

The current work has introduced the UMIA as new design conceived to minimize temporal bias when generating AI-designed horological images the algorithm is inspired by the unconscious mode of information processing by humans and can permit implementation through multiple bias-correcting schemes working side by side with main image construction the current technique openly addressed the widespread AI models' fallback to a 10:10 configuration when producing

watch faces and thereby demonstrated that this form of bias is possible to overcome with architectural innovation and not with data set selection alone.

The UMIA decreased by more than 76% the occurrence of 10:10 bias and hence effectively reestablished diversity for temporal representations of synthesized views of watches and system attained an average temporal display fidelity of roughly 0.91 on a 0–1 scale, whereas a baseline diffusion model yielded an average fidelity of roughly 0.63, corresponding to a 43.8% improvement based purely on temporal fidelity.

While the blending with stable diffusion showed there was neither loss of image quality nor any loss with its application through UMIA images possessed excellent realism and seemed to exhibit superior clarity through the application of UMIA transformations and the fact that end-users are always incapable of telling the current method outputs and baseline outputs apart other than at the correct representation of time is evidence of how subtle and effective the algorithm is.

This current work opens up numerous starting points for future work that can develop a general "unconscious bias mitigator" for generative models with broader applicability beyond a certain attribute or possibly direction is to develop a collection of micro-modules to cope with specific biases, for example, time bias, representation gender bias, object positioning bias, and so forth, to add upon request.

Moreover, using the UMIA paradigm to come up with text generation in correcting mistakes at a factual level or sound generation in preventing, e.g., a speech system from defaulting to some built-in accent would be a worthwhile thing to try out, and the initial assumption would always remain unchanged: identifying the bias, quantifying its effect, and incorporating a side-path of processing to neutralize its effect.

STATEMENT ON THE USE OF GENERATIVE ARTIFICIAL INTELLIGENCE

The images included in Figure 7-11 were generated using a custom Python-based image generation script, AdvancedUMIADrawer. The script used Python with Matplotlib, NumPy, pandas, and Pillow to generate watch images showing different times and visual styles. These images were generated for illustrative and research demonstration purposes and were incorporated into the manuscript as figure materials.

REFERENCES

- [1] Kang, M., Won, D., Luna, M., et al. (2023). Content preserving image translation with texture co-occurrence and spatial self-similarity for texture debiasing and domain adaptation. *Neural Networks*, 166: 722-737. <https://doi.org/10.1016/j.neunet.2023.07.049>
- [2] Wang, J., Liu, X., Di, Z., Liu, Y., Wang, X. (2023). T2IAT: Measuring valence and stereotypical biases in text-to-image generation. In *Findings of the Association for Computational Linguistics: ACL 2023*, Toronto, Canada, pp. 2560-2574. <https://doi.org/10.18653/v1/2023.findings-acl.160>
- [3] Safar, M., Omar, F.S., Safar, D., Jaafar, S., Mohammed, N.F. (2025). Unconscious mind-inspired algorithm: A novel approach to machine learning. *Ingénierie des*

- Systèmes d'Information, 30(3): 813-821. <https://doi.org/10.18280/isi.300324>
- [4] Shen, X., Plested, J., Caldwell, S., Gedeon, T. (2021). Exploring biases and prejudice of facial synthesis via semantic latent space. In 2021 International Joint Conference on Neural Networks (IJCNN), Shenzhen, China, pp. 1-8. <https://doi.org/10.1109/IJCNN52387.2021.9534287>
- [5] Karim, A.A., Lützenkirchen, B., Khedr, E., Khalil, R. (2017). Why is 10 past 10 the default setting for clocks and watches in advertisements? A psychological experiment. *Frontiers in Psychology*, 8: 1410. <https://doi.org/10.3389/fpsyg.2017.01410>
- [6] Safar, D., Safar, M., Jaafar, S., Al-Yachli, B.A., Shukri, A.K., Rasheed, M.H. (2025). Hallucinations in GPT-2 trained model. *Ingénierie des Systèmes d'Information*, 30(1): 31-41. <https://doi.org/10.18280/isi.300104>
- [7] Zhu, R.B., Liu, L., Ma, M.D., Li, H.X. (2020). Cognitive-inspired computing: Advances and novel applications. *Future Generation Computer Systems*, 109: 706-709. <https://doi.org/10.1016/j.future.2020.03.017>
- [8] Barbano, C.A., Tartaglione, E., Grangetto, M. (2025). Unsupervised learning of unbiased visual representations. *IEEE Transactions on Artificial Intelligence*, 6(5): 1171-1183. <https://doi.org/10.1109/TAI.2024.3514554>
- [9] Luccioni, S., Akiki, C., Mitchell, M., Jernite, Y. (2023). Stable bias: Evaluating societal representations in diffusion models. *Advances in Neural Information Processing Systems*, 36: 56338-56351. <https://doi.org/10.52202/075280-2458>
- [10] Tay, E., Lorenzi, M., Durrleman, S. (2025). Shape modeling of longitudinal medical images: From diffeomorphic metric mapping to deep learning. *Frontiers in Artificial Intelligence*, 8: 1671099. <https://doi.org/10.3389/frai.2025.1671099>
- [11] Xie, J.W., Zhu, S.C., Wu, Y.N. (2017). Synthesizing dynamic patterns by spatial-temporal generative ConvNet. In 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Honolulu, HI, USA, 1061-1069. <https://doi.org/10.1109/CVPR.2017.119>
- [12] Macmillan-Scott, O., Musolesi, M. (2025). (Ir)rationality in AI: State of the art, research challenges and open questions. *Artificial Intelligence Review*, 58: 352. <https://doi.org/10.1007/s10462-025-11341-4>
- [13] Friedrich, F., Brack, M., Struppek, L., et al. (2025). Auditing and instructing text-to-image generation models on fairness. *AI and Ethics*, 5(3): 2103-2123. <https://doi.org/10.1007/s43681-024-00531-5>
- [14] Elata, N., Kavar, B., Michaeli, T., Elad, M. (2024). Nested diffusion processes for anytime image generation. In Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV), Waikoloa, HI, USA, pp. 5018-5027. <https://doi.org/10.1109/WACV57701.2024.00493>
- [15] Chen, S.X., Vaxman, Y., Ben Baruch, E., et al. (2024). TiNO-Edit: Timestep and noise optimization for robust diffusion-based image editing. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, Seattle, WA, USA, pp. 6337-6346. <https://doi.org/10.1109/CVPR52733.2024.00606>
- [16] Vicente, L., Matute, H. (2023). Humans inherit artificial intelligence biases. *Scientific Reports*, 13: 15737. <https://doi.org/10.1038/s41598-023-42384-8>
- [17] Ge, H., Bastani, H., Bastani, O. (2024). Rethinking fairness for human-AI collaboration. *Leibniz International Proceedings in Informatics*, 287: 52:1-52:21. <https://doi.org/10.4230/LIPIcs.ITCS.2024.52>
- [18] Shen, T., Li, S., Wang, X.S., et al. (2024). Deep reinforcement learning enables better bias control in benchmark for virtual screening. *Computers in Biology and Medicine*, 171: 108165. <https://doi.org/10.1016/j.compbiomed.2024.108165>
- [19] Chen, M., Mao, S.W., Liu, Y.H. (2014). Big data: A survey. *Mobile Networks and Applications*, 19: 171-209. <https://doi.org/10.1007/s11036-013-0489-0>
- [20] Dhariwal, P., Nichol, A. (2021). Diffusion models beat GANs on image synthesis. *Advances in Neural Information Processing Systems*, 34: 8780-8794. <https://dl.acm.org/doi/10.5555/3540261.3540933>
- [21] Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, 54(6): 115. <https://doi.org/10.1145/3457607>
- [22] Newell, S., Marabelli, M. (2015). Strategic opportunities and challenges of algorithmic decision-making: A call for action on the long-term societal effects of datification. *Journal of Strategic Information Systems*, 24(1): 3-14. <https://doi.org/10.1016/j.jsis.2015.02.001>
- [23] Wei, X.H., Kumar, N., Zhang, H. (2025). Addressing bias in generative AI: Challenges and research opportunities in information management. *Information & Management*, 62(2): 104103. <https://doi.org/10.1016/j.im.2025.104103>
- [24] Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B. (2022). High-resolution image synthesis with latent diffusion models. In 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), New Orleans, LA, USA, pp. 10674-10685. <https://doi.org/10.1109/CVPR52688.2022.01042>