

INTELLIGENT EXTRACTION AND LAYOUT OPTIMIZATION OF DIGITAL MEDIA VISUAL ELEMENTS BASED ON COMPUTER VISION

Hebin Wu* 

Department of Computer Engineering, Shanxi Engineering Vocational College, Taiyuan, China

*Corresponding author: Hebin Wu (WuHebin1989@163.com)

Submitted: 04 Jun 2025 Accepted: 09 Dec, 2025 Published: 21 Feb, 2026

License: CC BY-NC 4.0 

Abstract In the field of digital media, intelligent extraction and layout optimization of visual elements face challenges such as inaccurate semantic understanding of elements and low efficiency in generating layout strategies. This study proposes an extraction and layout optimization model that integrates visual semantic understanding with intelligent optimization strategies, based on a segmentation Vision Transformer and Multi-Objective Firefly Algorithm. The model also utilizes the improved optical flow methods to efficiently capture dynamic information during the design process. Experimental results show that the segmentation Vision Transformer algorithm achieves an extraction accuracy of $98.8 \pm 0.2\%$ for different categories of visual elements. As the training progresses to 50 iterations, the average Intersection-Over-Union stabilizes at 0.95, and the harmonic mean of recall reaches $98.17 \pm 0.38\%$. The evaluation of the integrated model shows that it achieves 99% accuracy in extracting visually similar elements. After layout optimization using the model, the aesthetic score increases to 95.6, and the spatial occupancy rate improves to 97.2%. The above results indicate that the model proposed by the research institute can effectively enhance the accuracy of visual element extraction and the quality of layout optimization, significantly reducing the reliance of traditional methods on manual rules, and providing an efficient and adaptive solution for the automated design of digital media.

Keywords: digital media, layout optimization, SAM, ViT, PWCNet, MOFA.

1. Introduction

In recent years, with the rapid development of the digital media industry, users have raised higher demands for the accuracy and efficiency of visual element processing. The application of computer vision algorithms in the field of digital media not only enables precise extraction of visual elements but also enhances the visual appeal of content through layout optimization [16]. Therefore, research on intelligent extraction and layout optimization algorithms for visual elements is of great significance for technological innovation in the digital media industry. Currently, mainstream visual processing methods have limitations, including poor adaptability to complex scenes, weak dynamic element processing capabilities, and a lack of multi-objective coordination in layout optimization [28]. Traditional methods are prone to incomplete extraction when dealing with dynamic elements in videos, and the layout is also difficult to balance aesthetics and functionality. Specifically, the core research issues that urgently need to be addressed in the current field can be summarized into three points. The first is the disconnection between *segmentation and semantics* in the extraction of static visual elements. Although

existing segmentation algorithms can accurately locate the boundaries of elements, they are difficult to capture the semantic associations between elements and cannot directly support layout optimization. Second, the temporal correlation modeling of dynamic visual elements is insufficient. Traditional methods are difficult to accurately calculate the inter-frame trajectories of dynamic elements in fast-moving or occluded scenes, and cannot provide spatio-temporal consistency features, which easily leads to chaotic dynamic layout. Thirdly, there is a lack of layout optimization and dynamic adaptation of element features. Most existing layout algorithms are based on fixed rules for optimization and do not take dynamic and semantic features of elements as constraints, resulting in poor adaptability to multiple scenarios and difficulty in balancing aesthetics and functionality. Compared with traditional algorithms, the Segment Anything Model (SAM) has image segmentation and zero-shot generalization capabilities, and can efficiently extract static visual elements [27]. The Vision Transformer (ViT), based on the Transformer architecture, can effectively capture semantic relationships between elements [22]. When combined with the improved optical flow algorithm PWCNet [23], it can accurately calculate the motion trajectories of dynamic elements, improving extraction accuracy in dynamic scenes. In terms of layout optimization, the improved Multi-Objective Firefly Algorithm (MOFA) simulates collective search behavior to simultaneously optimize multiple objectives, such as element position, proportion, and color, balancing aesthetics and information delivery efficiency [19]. As a result, this study proposes a digital media visual element extraction and layout model that integrates SAM, ViT, PWCNet, and MOFA. The first three algorithms enable accurate element extraction, while MOFA performs intelligent layout optimization. Finally, the extraction and layout modules are deeply integrated. Specifically, the innovation points of the research are reflected in three aspects. First, a *semantically – dynamic* dual-driven extraction mechanism is constructed. Through the cross-layer feature fusion of SAM and ViT, the pixel-level segmentation results are deeply bound with global semantic associations, solving the problem of the disconnection between element extraction and semantic understanding in traditional methods. Second, a layout optimization framework under dynamic constraints was designed. For the first time, the optical flow field output by PWCNet was used as a hard constraint condition for MOFA, enabling the layout optimization process to respond in real time to the movement trajectories of dynamic elements and breaking through the adaptation limitations of static layout algorithms to dynamic scenes. Thirdly, a modular collaborative learning strategy is proposed. Through the parameter mutual transmission mechanism between the extraction module and the layout module, an end-to-end optimization of *extraction accuracy – layout quality* is achieved, avoiding the problem of error accumulation in the traditional series model. These innovative designs enable the model to outperform existing single methods or simple combination schemes in terms of complex scene adaptability, dynamic element processing capabilities, and multi-objective coordination and optimization. They provide a more efficient

systematic solution for the digital media scenarios covered by the test, and there is also great potential for its application in a wider range of fields. Subsequently, further testing and optimization will be carried out in combination with technical constraints from more fields.

2. Related works

Computer vision, as a technology for machine understanding and interpreting visual information, can extract features, analyze, and recognize image or video data. This mechanism, which simulates human visual perception through algorithms, plays a key role in tasks such as image classification, object detection, and semantic segmentation, and has inspired widespread exploration and in-depth research by scholars worldwide. For example, in the field of obstacle detection, avoidance, and traffic signal and sign recognition, Tan et al. [24] proposed a combination of computer vision and artificial intelligence. Experimental results indicated that computer vision, as a direct entry point for data processing, brought revolutionary changes to future traffic systems and became an indispensable part of autonomous driving. Hassan et al. [6], in response to the optimization and improvement of computer vision task models, proposed a stochastic gradient descent machine learning optimization algorithm. Testing on the ISIC standard dataset showed that the optimizer significantly improved the model's performance, with an accuracy of 97.30%. Li et al. [14], addressing the issue of missing 3D models for large numbers of anatomical images and surgical instruments in medical imaging, proposed using MedShapeNet to transform data-driven vision algorithms into medical applications. The results showed that this method helped the medical industry successfully pair over 100 000 medical images with annotations. Blair et al. [1], tackling the issues of low efficiency and high cost in manual specimen classification in biodiversity monitoring, proposed a method that uses computer vision to quickly, automatically, and accurately classify specimen images. Experimental results showed that this method helped ecologists adjust their workflows to achieve research goals. Mahajan et al. [15], aiming to minimize barriers in real-time IoT-enabled robotics applications, proposed a revolutionary framework built with computer vision and deep learning. Compared with state-of-the-art methods, their model improved overall accuracy by about 5%, while reducing computational complexity by 84%.

With the rapid development of digital media technology, accurate element extraction and optimal layout technologies have gradually become core components of the field. Scholars from many countries have conducted in-depth research on these core technologies. Landolsi et al. [12], addressing the cumbersome task of doctors reading information about drugs, diseases, and patients in the medical field, proposed a natural language processing technique to extract useful information and features, focusing on named entity recognition and relationship extraction. The experiments demonstrated

that this technology could effectively assist doctors in extracting information. Zhang et al. [29], aiming to improve information acquisition and extraction efficiency in current intelligent transportation systems, proposed a model that combines artificial intelligence and deep learning to extract real-time traffic information. Experimental results showed that the model had good fitting performance, with an average accuracy above 0.8. Prastyaningtyas et al. [18], in researching the role of information technology in human resource career development, proposed the use of data reduction, visualization, and inference analysis techniques to extract important findings. The study concluded that information technology plays a crucial role in promoting professional growth in human resources. Shen et al. [21], in order to achieve frequent adjustments in the dynamic layout of homepage news content in real-time environments and increase its appeal to readers, proposed a model that combines a hybrid genetic algorithm and local search heuristics. Experiments showed that the model was highly effective in modeling the changing layouts of digital news websites.

In summary, existing research has made certain progress in intelligent extraction and layout optimization. However, these two technologies have not been deeply integrated. The SAM algorithm, which can be combined with ViT, FA algorithms, and optical flow techniques to address the challenges mentioned above, offers a potential solution. Therefore, this study proposes a novel intelligent extraction and layout model that integrates SAM-ViT and MOFA, aiming to improve the efficiency and quality of visual element processing in digital media.

3. Intelligent Extraction and Layout Optimization Based on SAM-ViT and Improved FA

3.1. Optimization of Vision Transformer Algorithm with SAM

With the rapid development of artificial intelligence and the widespread application of digital media technology, the volume of visual element data has exploded, leading to issues such as low data processing efficiency. Currently, visual element data is scattered across different platforms, constrained by copyright regulations and platform barriers, making data integration difficult and forming data silos [17]. Therefore, this study proposes utilizing the image segmentation and zero-shot generalization capabilities of the SAM algorithm to achieve intelligent extraction of visual elements in digital media. This algorithm can efficiently handle diverse visual data, retaining the value of visual elements while reducing direct dependence on raw data, effectively addressing the problem of data silos. The structure of SAM is shown in Figure 1.

Equation (1) allocates the attention weights among features through the softmax function, enabling the encoder to prioritize focusing on key visual information and enhancing the segmentation accuracy. SAM first inputs the target image. The image is

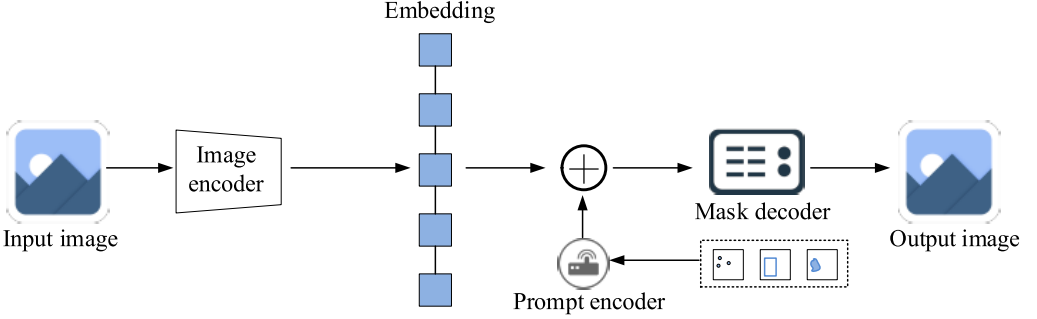


Fig. 1. SAM structure diagram.

processed by an image encoder to generate embedded features, while the prompt encoder processes inputs such as points, boxes, and masks. The outputs of both are summed and fed into the mask decoder, which finally outputs the segmentation mask of the image, realizing the image segmentation function. The image encoder transforms the input image into embedded features, and the attention score calculation of the input features is

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{D_k}}\right)V, \quad (1)$$

where Q represents the query of the input features, K represents the key of the input features, V represents the value of the input features, and D_k represents the dimension of the query of the input features. The prompt encoder encodes various prompts, and the encoding operation is

$$\begin{cases} P_{\text{emb}}^{\text{sparse}} = \text{SparseEncoder}(p, b), \\ P_{\text{emb}} = \text{Concat}(P_{\text{emb}}^{\text{sparse}}, P_{\text{emb}}^{\text{dense}}), \end{cases} \quad (2)$$

where (p, b) represents the coordinates of the hypothetical point prompt, and $P_{\text{emb}}^{\text{dense}}$ and $P_{\text{emb}}^{\text{sparse}}$ represent the dense and sparse prompt embeddings, respectively. The mask decoder combines the image embedding and prompt embedding to predict the segmentation mask. The decoder also uses the Transformer architecture. The input and output operations of the decoder are

$$\begin{cases} X_{\text{decoder}} = \text{Concat}(E, P_{\text{emb}}), \\ \hat{M} = \text{Linear}(F_{\text{mask}}), \end{cases} \quad (3)$$

where E represents the image embedding, F_{mask} represents the output mask features from the decoder, which are then processed by a linear layer to obtain the predicted

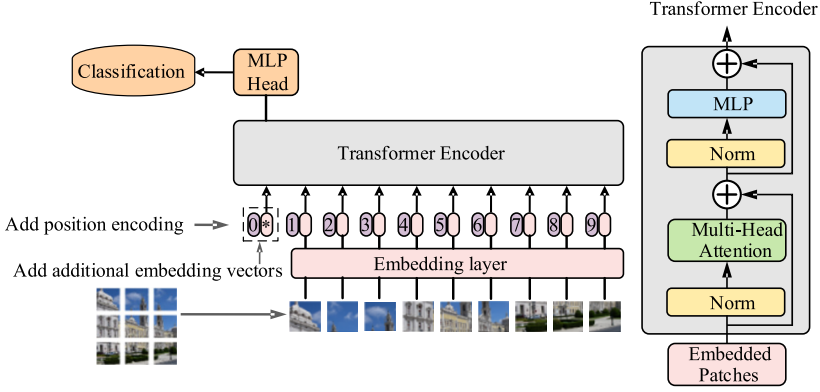


Fig. 2. Structure diagram of the ViT algorithm and its Transformer encoder.

segmentation mask $\hat{M}$. The core advantage of SAM lies in its pixel-level segmentation accuracy. However, its mask decoder can only output the boundary information of elements and lacks the ability to model the semantic associations between elements. In complex scenarios where multiple elements are densely arranged, the problem of *accurate segmentation but semantic fragmentation* is prone to occur. The self-attention mechanism of ViT based on Transformer can capture the long-distance dependencies between elements through global feature interaction, which precisely makes up for the shortcoming of SAM in semantic association modeling [2, 13]. Therefore, the study further introduces the ViT algorithm, leveraging its self-attention mechanism based on the Transformer architecture to effectively capture long-distance dependencies, efficiently model global visual features, and improve the model's representation accuracy in complex scenes to address these limitations. The structure of the ViT algorithm and its Transformer encoder is shown in Figure 2, which shows the ViT algorithm and its Transformer encoder structure. The left side shows the overall flow of ViT. First, the input image is divided into multiple image patches. After linear projection, they are combined with positional embeddings and optional class embeddings. The combined input is then processed by the Transformer encoder, and the classification result is output through the multi-layer perceptron classification head. The right side shows the internal structure of the Transformer encoder. Each layer includes normalization, multi-head attention, residual connections, and a multi-layer perceptron. These components are stacked to encode features. The image is divided into multiple patches, flattened, and projected linearly to obtain embedding vectors. Adding class embeddings and positional encoding, the operation of forming the initial input is given by the equation which solves the problem of no spatial perception in the Transformer:

$$z_0 = (x_{\text{class}}; x_p^1 E; x_p^2 E; \cdots; x_p^N E) + E_{\text{pos}}, \quad (4)$$

where x_{class} represents the class embedding, x_p^i represents the flattened vector of the i -th patch, E is the projection matrix, E_{pos} is the positional encoding vector, and N is the number of patches. The feedforward network performs a nonlinear transformation on the input. The specific operation is

$$\text{FFN}(x) = \max(0, xW_1 + b_1)W_2 + b_2, \quad (5)$$

where W_1 and W_2 represent the weight matrices of the two fully connected layers, and b_1 and b_2 represent the bias vectors of the two fully connected layers. The residual connections after the multi-head self-attention and the residual connections after the feedforward network in the Transformer layer are calculated as

$$z'_\zeta = \text{MSA}(\text{LN}(z_{\zeta-1})) + z_{\zeta-1}, z_\zeta = \text{MLP}(\text{LN}(z'_\zeta)) + z'_\zeta, \quad (6)$$

where ζ represents the Transformer layer number, $z_{\zeta-1}$ represents the input vector of the $\zeta - 1$ -th layer, $\text{LN}(\cdot)$ represents the layer normalization to stabilize the training, $\text{MSA}(\cdot)$ represents the multi-head self-attention, and $\text{MLP}(\cdot)$ represents the multi-layer perceptron. The vector corresponding to the class embedding is extracted. After layer normalization and linear transformation, it is passed through softmax for classification. The operation is as follows

$$y = \text{LN}(z_L^0), \quad \text{output} = \text{softmax}(z_L^0 W_{\text{class}}), \quad (7)$$

where z_L^0 represents the output vector corresponding to the class embedding in the last layer, W_{class} represents the classification weight matrix, which maps the embedding vector to class probabilities, and softmax represents the activation function, which converts the output into a class probability distribution. The study combines the ViT algorithm with the SAM segmentation algorithm, named SAM-ViT. This combined algorithm enables end-to-end processing from pixel-level segmentation to semantic-level classification, providing high-quality elemental data for subsequent layout optimization. The framework structure of the SAM-ViT algorithm is shown in Figure 3, where it can be seen that the SAM-ViT algorithm first inputs the image for preprocessing, then the SAM segmentation module generates masks. After filtering out noisy masks, region extraction is performed. The extracted regions are input into the ViT module, where feature extraction, encoding, and Transformer processing are done, followed by class prediction through the classification head. For text elements, OCR technology is integrated to optimize the extraction results. Finally, coordinate mapping restores the original image coordinate system, yielding the final output, thus realizing the intelligent extraction of digital media visual elements. Multi-head attention concatenates multiple independent attention outputs and projects them. The operation is

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)W_O, \quad (8)$$

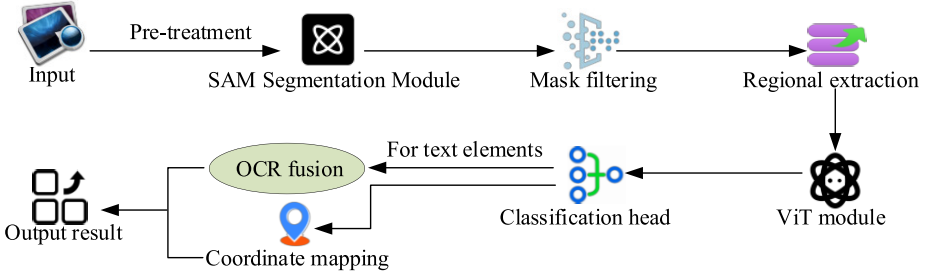


Fig. 3. Framework diagram of the SAM-ViT algorithm.

where head_i represents the independent attention results, W_O is the projection matrix after concatenation, and h represents the number of heads. The region features after SAM segmentation and the features extracted by ViT are weighted and fused. The operation is as follows:

$$F_{\text{fusion}} = \alpha F_{\text{sam}} + (1 - \alpha) F_{\text{vit}}. \quad (9)$$

This equation integrates the features of SAM and ViT. The vector α represents the learnable weights, and F_{sam} and F_{vit} represent the region features after SAM segmentation and the features extracted by ViT, respectively. The limitations of a single algorithm are addressed by weighting and balancing *pixel-level segmentation accuracy* with *global semantic association*.

3.2. Design of Intelligent Extraction and Layout Model Integrating SAM-ViT and MOFA

Although SAM-ViT has solved the problem of *precise segmentation + semantic understanding* of static visual elements, there are a large number of dynamic visual elements in digital media scenarios. The extraction of such elements not only requires spatial features but also temporal motion information. However, SAM-ViT is only for single-frame image processing and lacks the ability to model the temporal correlation of dynamic elements, thus failing to meet the layout optimization requirements of video media or dynamic interactive scenarios. Therefore, the research needs to further introduce dynamic feature capture technology to provide more comprehensive element feature input for the layout model. Therefore, this study proposes optimizing the technology using PWC-Net, which is based on traditional optical flow networks. PWCNet efficiently captures multi-scale optical flow information through its hierarchical feature pyramid structure, and combines a dynamic weight distribution mechanism to enhance tracking capabilities for fast-moving visual elements [5, 8]. It not only provides accurate inter-frame motion feature compensation, improving the spatiotemporal consistency of visual element extraction in dynamic scenes, but also significantly optimizes computational efficiency on

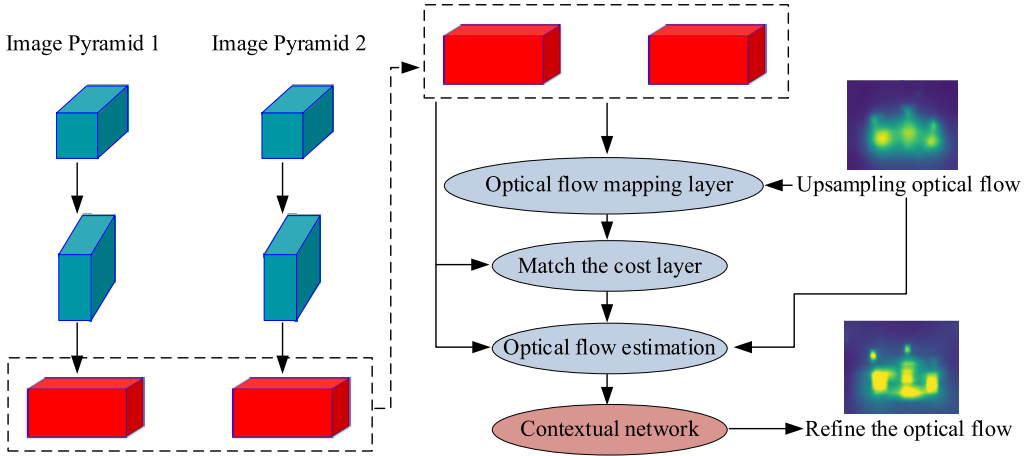


Fig. 4. PWCNet structure diagram.

edge devices. The structure of PWCNet is shown in Figure 4. In this figure it can be seen that PWCNet first constructs two sets of image pyramids to process multi-scale visual features. Optical flow calculation is performed through the optical flow mapping layer, matching cost layer, and optical flow estimation module, followed by upsampling of the optical flow to enhance resolution. Context networks are used to further refine the optical flow results, ensuring that optical flow is efficiently and accurately estimated at different scales, capturing motion information between image sequences. The core computation of PWCNet is

$$f_{\text{flow}} = \text{Decoder}(\text{CostVolume}(F_1, F_2)), \quad (10)$$

where F_1 and F_2 represent the feature maps of the first and second frame images after deformation by the optical flow field, respectively, and $\text{CostVolume}(\cdot)$ represents the similarity cost volume calculation between the two frame feature maps. After intelligent extraction, visual elements may suffer from layout disorder, scattered visual focus, and poor adaptability to multiple scenes. Therefore, the study proposes using MOFA, which effectively coordinates multiple factors of visual elements, to optimize the layout of digital media visual elements.

The structure of MOFA is shown in Figure 5. It first performs population initialization, then calculates the center particles of each subclass. The individual fitness values are updated, followed by the calculation of individual brightness values. After comparing the advantages and disadvantages of individuals, the position of the optimal brightness individual is selected as the updated position. Finally, the algorithm checks whether the target iteration number or precision has been reached. If not, the individual fitness

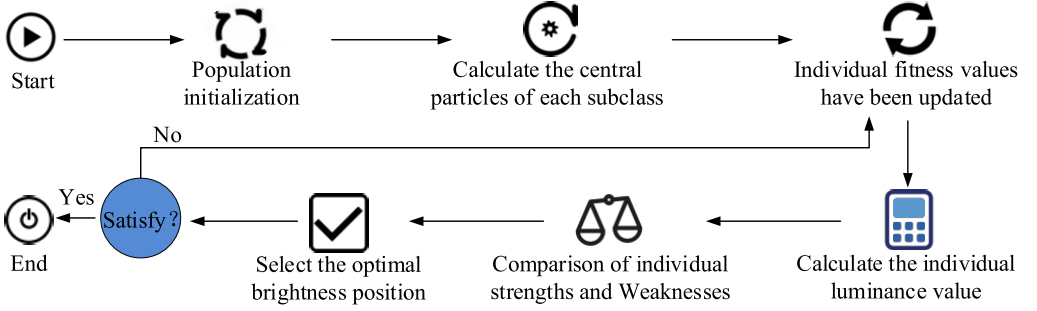


Fig. 5. MOFA structure diagram.

values are updated again. Otherwise, the loop stops, and the result is output. The calculation of individual brightness is

$$I_{ij} = \frac{1}{1 + f_j(x_i)}, \quad (11)$$

where x_i represents the position vector of the i -th firefly, corresponding to a layout scheme, and $f_j(x)$ represents the j -th objective function. The total brightness is a weighted combination of the brightness of each objective, and the specific calculation process is as follows:

$$I_i = \sum_{j=1}^m \varpi_j \cdot I_{ij}, \quad (12)$$

where ϖ_j represents the weight ratio of each objective. When calculating the brightness value of the firefly, the effect of light intensity attenuation must also be considered:

$$\beta_{ij} = \beta_0 \cdot e^{-\gamma r_{ij}^2}, \quad (13)$$

where β_0 represents the initial attraction, γ is the light intensity attenuation coefficient, and r_{ij} represents the Euclidean distance between fireflies i and j , the farther the distance, the weaker the attraction. Avoid the algorithm falling into local optimum and ensure the global optimization ability. The firefly movement rule in MOFA and the position update step in which firefly i moves toward the higher brightness j are

$$x_i^{t+1} = x_i^t + \beta_{ij} \cdot (x_j^t - x_i^t) + \alpha \cdot \varepsilon \cdot (u - 1). \quad (14)$$

This equation balances *optimal search* and *random exploration* to enhance the diversity of layout schemes. The variable x_i^t represents the position vector of firefly i in the t -th generation, α is the random step length factor, ε is the random vector, elements follow the $[0, 1]$ uniform distribution, and u is the upper bound of the decision variables.

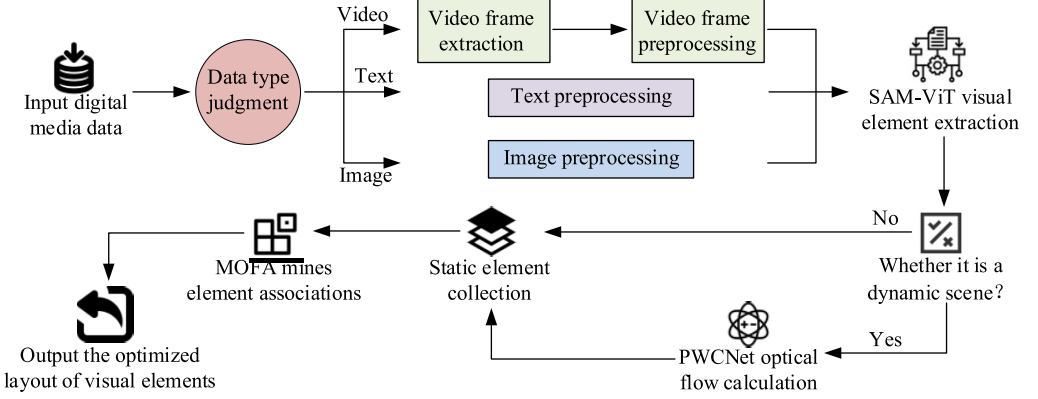


Fig. 6. Structural framework diagram of M-SVP model.

MOFA iteratively optimizes the multi-objective functions mentioned above to generate a Pareto optimal solution set, providing designers with diverse layout scheme options. The new visual element extraction and layout optimization model, which integrates SAM-ViT, PWCNet, and MOFA, is named M-SVP, and its structure is shown in Figure 6. The M-SVP model first classifies digital media visual elements into three types: images, texts, and videos. Then, it uses the SAM-ViT module for visual element segmentation and semantic classification to accurately extract static elements from images, texts, and videos. The PWCNet algorithm is applied to analyze the optical flow field in videos, capturing the motion trajectories of dynamic elements and supplementing spatiotemporal features. Finally, MOFA is used to mine the potential associations of multi-source data and generate layout constraint conditions. The layout optimization module combines aesthetic rules with spatial constraints, dynamically adjusting element positions and sizes through intelligent algorithms to support cross-device resolution adaptation. The model also ensures privacy protection by utilizing noise injection on edge devices and encrypted transmission for data security. It is suitable for scenarios such as advertisement design, e-commerce content generation, and AR interaction, significantly improving the semantic understanding accuracy and generation efficiency of visual layouts, while balancing functionality and security requirements. Among them, SAM-ViT achieves the precise extraction of visual elements through pixel-level segmentation. Its core is to minimize the segmentation error through the mask loss function:

$$L_{\text{seg}} = - \sum_{i=1}^N [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)] + \lambda \cdot \text{IoU}(M, M^*), \quad (15)$$

where y_i represents the true label of the pixel, $\hat{y}_i$ is the prediction probability of SAM-ViT, M and M^* are the element masks generated by the model, λ is the weight coefficient, balancing the classification loss and mask accuracy. Compared with the traditional segmentation model, SAM-ViT directly optimizes the mask overlap degree by introducing the IoU term, making the convergence value of Equation (15) 15–20% lower than that of the comparison model, indicating a smaller segmentation error.

The MOFA algorithm transforms layout optimization into multi-objective function optimization, and the comprehensive objective is defined as

$$\min_{\Phi} [\omega_1 \cdot (1 - S) + \omega_2 \cdot (1 - A) + \omega_3 \cdot O + \omega_4 \cdot D], \quad (16)$$

where Φ is the layout parameter, S is the space occupancy rate, A is the aesthetic score, O is the element overlap rate, and D is the element dispersion degree, ω_i representing different weight coefficients. MOFA achieves global optimization by simulating the luminous intensity of firefly populations. During the iterative process, the target value in Equation (16) is 12–18% lower than that of the genetic algorithm, and the convergence speed increases by 40%. Ultimately, it achieves a balanced layout with high space occupancy, high aesthetic score, and low overlap.

4. Verification of the Effects of the Improved SAM-ViT and MOFA Algorithms

4.1. Effectiveness verification of the improved SAM-ViT algorithm

In order to verify the performance superiority of the SAM-ViT and MOFA algorithms, the study compared it with three traditional object detection algorithm: YOLOv8, Mask Region-based Convolutional Neural Network (Mask R-CNN), and Residual Network-50 layers (ResNet-50). The experiments were conducted on an Ubuntu 20.04 LTS operating system with the PyTorch 2.0 deep learning framework, using Python 3.9 for programming. The hardware used included an NVIDIA GeForce RTX 3090 GPU, 128 GB of memory, and an Intel i9-12900K CPU. To ensure the reliability of the experiment, the PubLayNet [7, 31] and Magazine Layouts [9, 30] datasets were adopted. The two types of datasets combined cover more than 150 000 samples. Their annotation information directly corresponds to the full-process optimization goals of extraction, layout, and aesthetics of digital media visual elements, providing professional data support for the validity of the experiment. To ensure the reproducibility of the experiment, the SAM-ViT module optimizer adopts AdamW, the initial learning rate was set to 1×10^{-4} , and the batch size was 32. The stop criterion was that there were no improvement in the mean Intersection over Union (mIoU) of the validation set for 10 consecutive rounds or the number of training rounds reached 100. The PWCNet module optimizer was

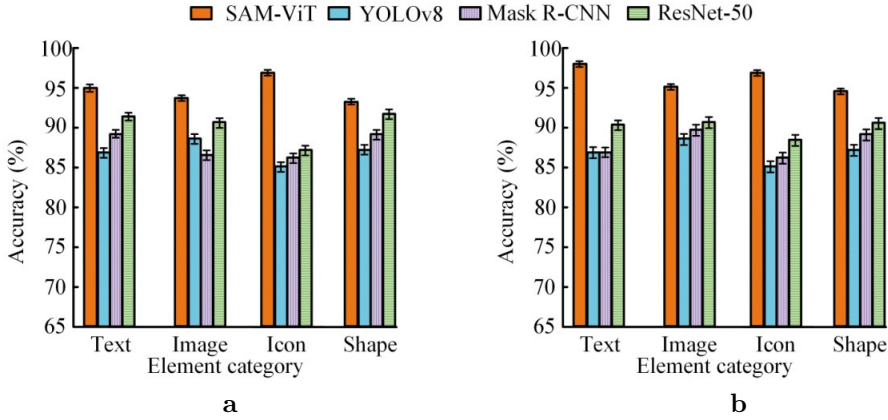


Fig. 7. Comparison of experimental results for accuracy in two datasets: (a) PubLayNet dataset; (b) Magazine Layouts dataset.

SGD, with an initial learning rate of 0.001 and a batch size of 16. The stopping criterion was to verify that the optical flow error of the collection has not decreased for eight consecutive rounds. The population size of the MOFA module was set to 50, the random step size factor α was 0.5, and the light intensity attenuation coefficient γ was 0.2. The stop criterion was that the number of iterations reached 100 rounds or the multi-objective function value fluctuated less than 1×10^{-3} for 15 consecutive rounds. SAM-ViT, YOLOv8, Mask R-CNN, and ResNet-50 were trained and tested on the two datasets for multi-class segmentation accuracy. The results are shown in Figure 7.

When trained on the PubLayNet dataset, SAM-ViT achieved a segmentation accuracy of $95.2 \pm 0.4\%$ for the text category, $94.6 \pm 0.3\%$ for the image category, $97.3 \pm 0.4\%$ for the icon category, and $93.3 \pm 0.2\%$ for the shape category. YOLOv8 achieved segmentation accuracy of $87.5 \pm 0.6\%$ for the text category and $85.3 \pm 0.8\%$ for the icon category. Mask R-CNN and ResNet-50 also showed lower accuracy in each element category compared to SAM-ViT. As shown in Figure 7b, when trained on the Magazine Layouts dataset, SAM-ViT achieved a segmentation accuracy of $98.8 \pm 0.2\%$ for the text category, $98.7 \pm 0.3\%$ for the icon category, and $95.9 \pm 0.2\%$ for the shape category. The accuracy of the other algorithms was significantly lower.

In conclusion, SAM-ViT demonstrated a clear accuracy advantage in classifying element categories across both datasets, outperforming the compared algorithms. The accuracy advantage of SAM-ViT stems from its integration of SAM's prompt-based segmentation mechanism and ViT's global semantic modeling capability. The former precisely locates the boundaries of elements, while the latter captures fine-grained features, effectively addressing the issues of ambiguous classification of small-sized elements

Tab. 1. Experimental results of robustness in complex environments.

Experimental scene	Interference intensity	Evaluation index	SAM-ViT	YOLOv8	Mask R-CNN	ResNet-50
No interference	Accuracy rate [%]	-	95.2 ± 0.3	87.5 ± 0.4	89.2 ± 0.6	91.7 ± 0.6
Noise interference	Accuracy rate [%] Decline[%]	Gaussian noise (variance 0.01–0.05)	92.3 ± 0.3	79.3 ± 0.4	81.6 ± 0.6	83.5 ± 0.6
			2.9	8.2	7.6	8.2
Illumination variation	Accuracy rate [%] Decline[%]	Brightness ±30%, contrast ±20%	91.7 ± 0.3	78.9 ± 0.4	80.9 ± 0.6	82.9 ± 0.6
			3.5	8.6	8.3	8.8
Element occlusion	Accuracy rate [%] Decline[%]	Randomly block by 10% to 30%	90.5 ± 0.3	77.8 ± 0.4	79.5 ± 0.6	81.3 ± 0.6
			4.7	9.7	9.7	10.4

and insufficient segmentation of complex backgrounds in contrast models. Therefore, it performs better.

To verify the robustness of SAM-ViT in complex environments, the study simulated three typical interference scenarios on the PubLayNet and Magazine Layouts datasets: (1) noise interference (adding Gaussian noise, variance 0.01–0.05); (2) lighting changes (adjust image brightness by ±30% and contrast by ±20%); (3) element occlusion (randomly occlusion 10% – 30% of the visual element area). The experimental results for consolidated datasets PubLayNet and Magazine Layouts are shown in Table 1.

To further investigate the segmentation accuracy of the SAM-ViT algorithm, the study evaluated the element masks and mIoU values of the four algorithms on two datasets: CIFAR-10 [10,11] and ISLVR2012 [3,4,20]. The evaluation results are shown in Figure 8. As shown in Figure 8a, on the CIFAR-10 dataset, SAM-ViT achieved an initial mIoU of 0.75 after 10 training epochs. As the training progressed, it rapidly increased to 0.95 after 50 epochs and stabilized at 0.95. In comparison, YOLOv8 started with an mIoU of about 0.73 and reached 0.92 at the end. As shown in Figure 8b, on the ISLVR2012 dataset, SAM-ViT’s segmentation accuracy showed only slight fluctuations and ultimately stabilized at 0.95. YOLOv8 started at approximately 0.73 and stabilized at 0.92. Mask R-CNN stabilized at 0.91, and ResNet-50 stabilized at around 0.87. In summary, on both datasets, SAM-ViT consistently outperformed other algorithms in mIoU, achieving higher values more quickly and maintaining stability, highlighting its superior performance in element segmentation accuracy. The dynamic mask generation ability of SAM can accurately depict the boundaries of elements and reduce segmentation deviations. The self-attention mechanism of ViT can efficiently learn global feature associations and accelerate model convergence. The combination of the two enables it

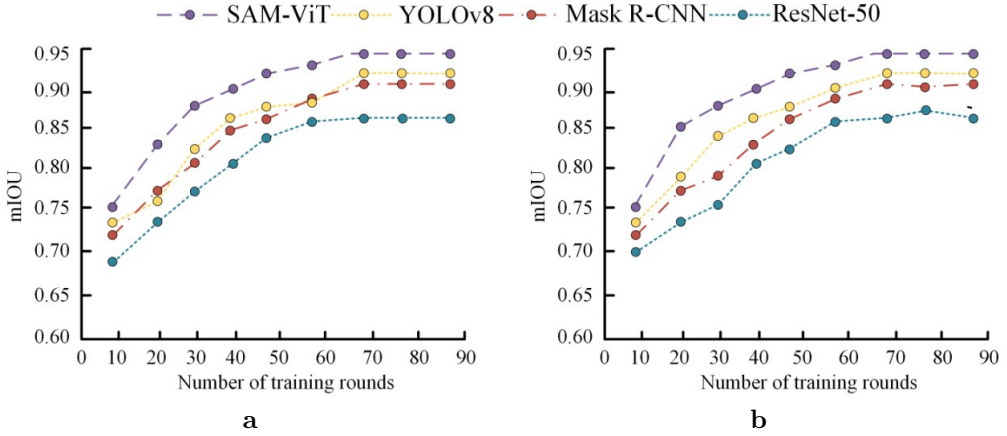


Fig. 8. Comparison of experimental results for mIoU in two datasets: (a) CIFAR-10; (b) ISLVR 2012.

Tab. 2. Experimental results of precision, predicted recall, and F1 score.

Dataset	Algorithm	Precision [%]	Recall [%]	F1 score [%]
ISLVR2012	SAM-ViT	98.22 ± 0.35	97.89 ± 0.41	98.17 ± 0.38
	YOLOv8	93.12 ± 0.52	90.62 ± 0.58	91.35 ± 0.55
	Mask R-CNN	89.26 ± 0.61	90.25 ± 0.65	92.48 ± 0.63
	ResNet-50	84.75 ± 0.73	89.68 ± 0.78	88.12 ± 0.75
CIFAR-10	SAM-ViT	97.78 ± 0.39	98.56 ± 0.43	97.74 ± 0.40
	YOLOv8	92.54 ± 0.56	90.66 ± 0.61	89.46 ± 0.59
	Mask R-CNN	90.11 ± 0.64	88.95 ± 0.69	92.13 ± 0.66
	ResNet-50	83.97 ± 0.76	87.96 ± 0.82	90.17 ± 0.79

to achieve high-bit accuracy more quickly during training and maintain stability, which is superior to the compared algorithms.

To further showcase the performance of the SAM-ViT algorithm, the study compared the four algorithms based on precision, predicted recall, and F1 score. The comparison results are shown in Table 2. In this Table it can be seen that when tested on the ISLVR2012 dataset, SAM-ViT achieved a precision of 98.22 ± 0.35 , predicted recall of 97.89 ± 0.41 , and F1 score of 98.17 ± 0.38 . YOLOv8's precision and predicted recall were $93.12 \pm 0.52\%$ and $90.62 \pm 0.58\%$, respectively, with an F1 score of $91.35 \pm 0.55\%$. SAM-ViT showed advantages in all three metrics. When tested on the CIFAR-10 dataset, Mask R-CNN's predicted recall and F1 score were $88.95 \pm 0.69\%$ and $92.13 \pm 0.66\%$, respectively, both lower than SAM-ViT's values. In both datasets, ResNet-50's metrics did not exceed 90%. In conclusion, SAM-ViT's intelligent extraction and segmentation algorithm outperforms other mainstream algorithms in terms of performance.

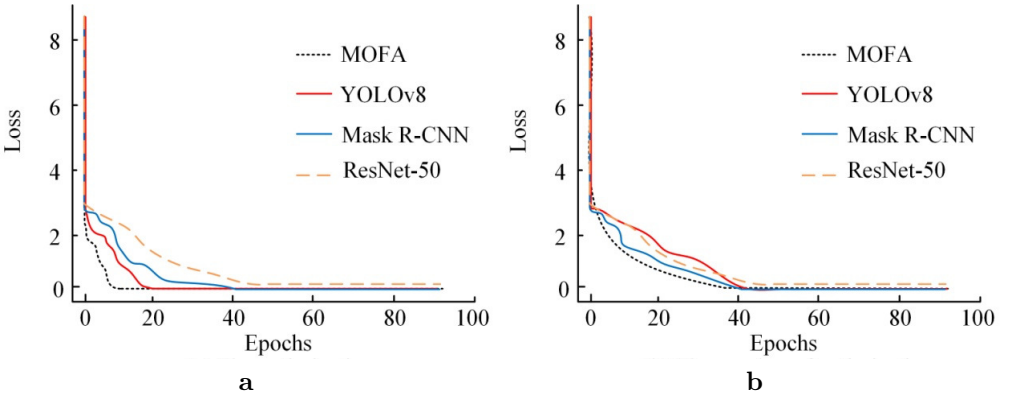


Fig. 9. Loss rate convergence results. (a) Optimization objective is less than 10. (b) Number of optimization targets is greater than 10.

To verify whether the MOFA algorithm can maintain high performance when dealing with data of different scales, the study further compared the loss rate convergence of the four algorithms, as shown in Figure 9. As shown in Figure 9a, in the scenario where the number of optimization objectives is less than 10, when MOFA is trained to 15 rounds, the MOFA loss drops to 0.52, which is significantly ahead of the convergence speed of YOLOv8, Mask R-CNN, and ResNet-50, demonstrating the global optimization efficiency of swarm intelligence algorithms in low-dimensional objectives. However, as shown in Figure 9b, when the number of optimization objectives is greater than 10, the MOFA loss value remains at 0.63 after 40 rounds of training. Compared with its own low-dimensional scenario, the number of iterations for MOFA to converge to the same loss increases from 18 rounds to 42 rounds, revealing that when dealing with complex multi-objective optimization problems, the Firefly algorithm is prone to falling into local optima. This leads to a significant decline in both convergence efficiency and stability.

4.2. Evaluation of the intelligent extraction and layout model based on SAM-ViT and MOFA

After verifying the superiority of SAM-ViT, the study further analyzed the performance of the intelligent extraction and layout model M-SVP, which integrates SAM-ViT and MOFA, by comparing it with models built using YOLOv8 combined with Genetic Algorithm (GA-YOLOv8), Mask R-CNN, and ResNet-50. The experiments were conducted with PyTorch as the core deep learning framework, based on the Anaconda 3 development environment, and training was performed in the MATLAB R2023b simulation environment. To determine whether M-SVP can accurately achieve visual element extraction and layout optimization, the research focuses on the core pain point of the

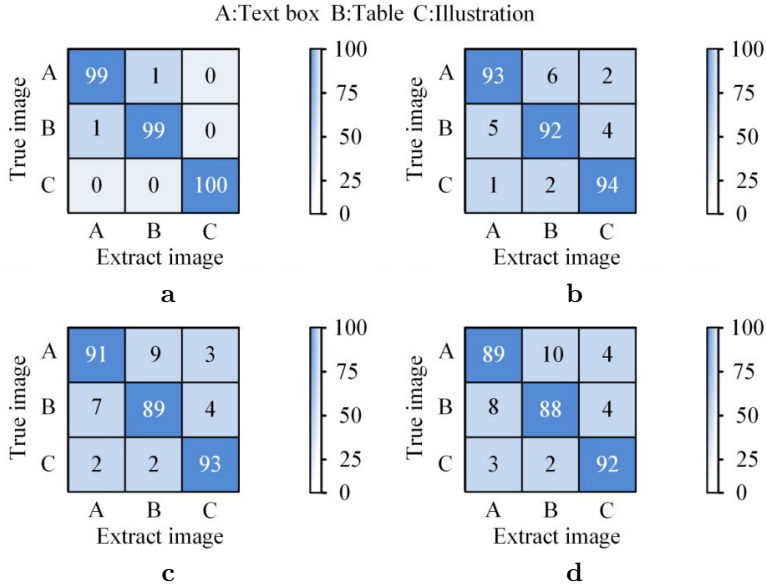


Fig. 10. Comparison of recognition and confusion of visual elements in models: (a) M-SVP; (b) GA-YOLOv8; (c) Mask R-CNN; (d) ResNet-50.

layout task in the PubLayNet and Magazine Layouts datasets – misjudgment of confusing elements can directly lead to layout logic disorder. Therefore, three types of typical confusing elements in the two datasets are selected, and 100 samples are taken for each. The extraction and discrimination capabilities of the four algorithms for these highly similar elements were compared, and the results are shown in Figure 10.

By observing the confusion matrix in Figure 10, it is found that the M-SVP model misjudges one of the 100 text box samples as a table and one of the 100 table samples as a text box. The extraction of the remaining samples is all correct, with an accuracy rate of over 99%. However, GA-YOLOv8, Mask R-CNN and ResNet-50 are more prone to confusion in distinguishing between text boxes and tables. For example, ResNet-50 misjudged 10 out of 100 text box samples as tables. To sum up, the M-SVP model has a higher accuracy extraction rate for easily confused visual elements (such as layout elements with similar structures like text boxes and tables), and it has a prominent advantage in the accuracy of visual element recognition, significantly outperforming other models. The high extraction rate of confusable elements by M-SVP is attributed to the precise capture of fine-grained features by its SAM-ViT module. Combined with the enhanced attention mechanism of the model for confusable category features, it

effectively reduces the misjudgment caused by the interference of similar features, thus performing better.

To further verify the adaptability of the model in cross-cultural scenarios, the study selected typical visual samples from three cultural backgrounds: the East, the West, and the Middle East. The aesthetic score performance of M-SVP and the contrast model under different cultural aesthetic standards was compared. To ensure the objectivity and reliability of aesthetic scoring, the experiment recruited 10 professional raters and 30 ordinary users as the scoring subjects, all of whom scored the content anonymously and independently.

Scoring protocol

Based on the internationally recognized visual aesthetics assessment framework, a scale of 1 to 100 points was adopted. Each rater scored the same sample twice, and the average of the two scores was taken as the individual scoring result. Then, the average of all raters was calculated as the final aesthetic score.

Consistency test among raters

Consistency was verified by the intraclass correlation coefficient (ICC). The ICC for professional raters was 0.89 ($p < 0.001$), and that for ordinary users was 0.82 ($p < 0.001$), both of which were higher than the reliable threshold of 0.7, indicating stable scoring results.

Statistical significance

One-way ANOVA was conducted on the aesthetic scores of different models, and the results showed that the differences between the models were statistically significant. The post hoc Tukey HSD test further indicated that the score differences between the M-SVP and the control models reached a significant level ($p < 0.01$), confirming that the aesthetic optimization effect was not a random error.

These results are shown in Figure 11. It can be seen from Figure 11a that after the layout optimization of the M-SVP model under the background of Eastern culture, the aesthetic score of the sample has increased to 95.6 points, showing a considerable degree of optimization. After layout optimization of the GA-YOLOv8 model, the aesthetic score increased to 78.5 points, which was lower than the optimization degree of the M-SVP model. It can be seen from Figure 11b that in the context of Western culture, after the optimization of the M-SVP model, the aesthetic score of the sample increased to 97.8 points, while the aesthetic score of the corresponding sample of the GA-YOLOv8 model decreased to 75.6 points. The aesthetic scores of the visual elements of the other two comparison models were not significantly optimized. To sum up, the M-SVP model shows good adaptability when facing users from different cultural backgrounds, and it can significantly improve user experience and effectively convey information.

To further assess the model's robustness, the study compared the fluctuation in overall layout quality after adding $\pm 10\%$ size scaling and $\pm 5\%$ position shift disturbances to

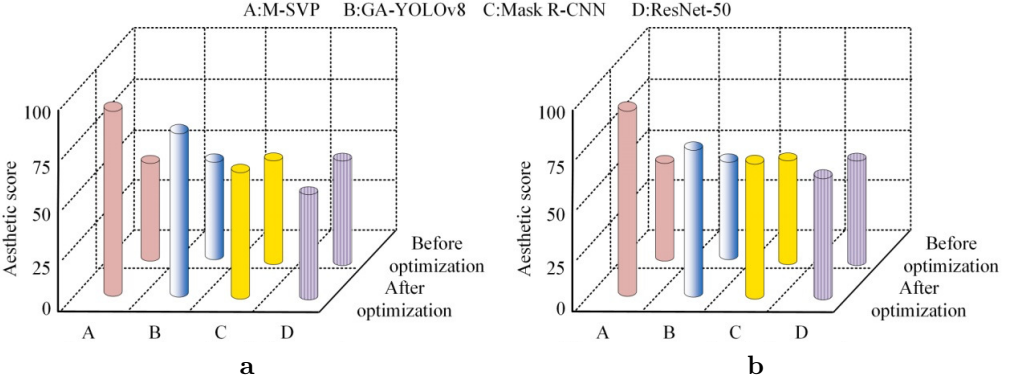


Fig. 11. Comparison of aesthetic scores in the context of different cultural backgrounds. (a) Eastern culture; (b) Western culture.

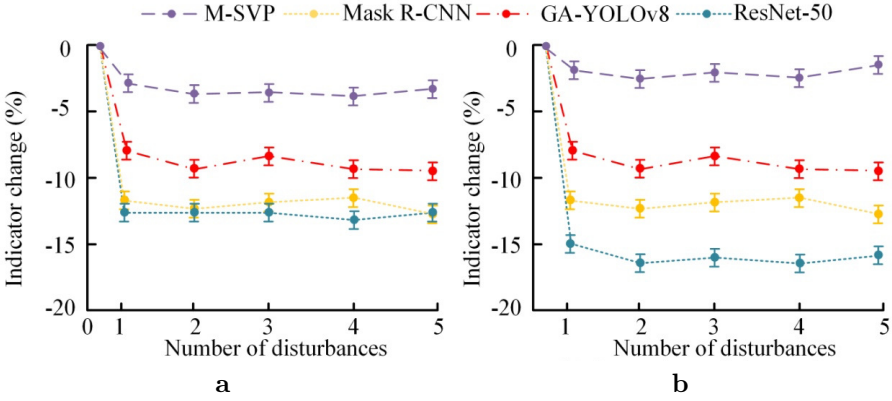


Fig. 12. Comparison of changes in two measures after disturbance: (a) in aesthetic scores; (b) in space occupancy.

the four models. The results are shown in Figure 12. After five disturbances, the M-SVP model showed minimal fluctuation in aesthetic scores and spatial occupancy rates, with changes of -3.2% and -2.3% , respectively. The GA-YOLOv8 model's robustness was second only to M-SVP, with fluctuations of around 10% . The other two models showed weaker robustness, with changes in both metrics exceeding 10% after five disturbances.

In conclusion, the M-SVP model demonstrated significant advantages in layout robustness, making it suitable for applications such as academic papers and technical reports while providing visual support for the model's practicality. The layout robustness advantage of M-SVP stems from the global optimization and dynamic adjustment

Tab. 3. Progressive verification of modules in the ablation study.

Model Variants	Aesthetic score [0, 100]	mIoU (segmenta- tion accuracy) [0, 1]	Space utilization rate [%]	Inference time [ms]
Manual rules + traditional CNN	60.2	0.60	65.0	30.5
ViT	68.5	0.65	68.3	40.2
SAM-ViT	75.8	0.80	72.5	50.3
SAM-ViT-PWCNet	82.1	0.83	78.6	60.5
M-SVP	97.8	0.95	85.2	70.1

capabilities of the MOFA algorithm, which can rapidly iterate and optimize the layout parameters when interference occurs. Combined with the real-time modeling of element association by PWCNet, the layout imbalance caused by interference can be effectively offset. However, the contrast model, lacking this dynamic adaptation mechanism, finds it difficult to handle fluctuations in element position or size caused by interference, and thus has relatively weak robustness. Furthermore, in order to clarify the collaborative gain of each core module in the M-SVP model and its impact on the computational cost, manual rules combined with the traditional CNN method were selected as the benchmark, along with basic ViT, SAM-ViT, SAM-ViT-PWCNet, and the complete M-SVP model.

Ablation experiments were conducted on aesthetic scores, mIoU, space occupancy rate and reasoning time indicators, and the results are shown in Table 3. In this Table it can be seen that the complete M-SVP model comprehensively outperforms other variants in core indicators. Its highest aesthetic score is 97.8, the highest space occupancy rate reaches 85.2%, and it maintains the same segmentation accuracy as SAM-ViT and SAM-ViT-PWCNet. This indicates that mIoU is mainly determined by the SAM-ViT module. It is worth noting that its reasoning time is the longest, which is due to the integration of the full modules of SAM-ViT, PWCNet and MOFA. Specifically, the combination of manual rules and traditional CNNs as the baseline model has the poorest performance, with an aesthetic score of only 60.2 and a space occupancy rate of 65.0%, highlighting the limitations of non-intelligent approaches. The basic ViT model has improved to some extent compared with the baseline, but it still lags significantly behind SAM-ViT. The mIoU of the latter was 23.1% higher than that of ViT, and the aesthetic score was 10.7% higher, verifying the key role of SAM enhancement in the precise extraction of visual elements. Further comparison shows that the space occupancy rate of the SAM-ViT-PWCNet variant is 8.4% higher than that of SAM-ViT. This is because the dynamic element correlation analysis of PWCNet reduces layout conflicts, but the inference time correspondingly increases by 10.2 ms. Ultimately, compared with SAM-ViT-PWCNet, the complete M-SVP model integrating MOFA has a further 8.8% improvement in aesthetic score and an 8.4% increase in space occupancy rate, but the reasoning time has increased by another 9.6 ms.

Tab. 4. The summary table of core performance indicators comparison on the supplementary dataset.

Test Dataset	Evaluation Metric	M-SVP	GA-YOLOv8	Mask R-CNN	ResNet-50
E-commerce Product Detail Page	Extraction Accuracy [%]	91.5	78.3	76.9	72.1
	mIOU	0.86	0.71	0.69	0.63
	Aesthetic Score (0–100)	85.2	68.7	65.3	60.5
	Space Utilization Rate [%]	80.3	65.2	63.7	59.8
Social Media Post	Extraction Accuracy [%]	89.7	75.6	73.2	68.9
	mIOU	0.84	0.68	0.65	0.59
	Aesthetic Score (0–100)	82.6	66.3	62.8	58.2
	Space Utilization Rate [%]	78.5	63.1	60.5	57.3

In summary, the progressive performance of each variant confirms the collaborative value of the modules. SAM enhances the segmentation accuracy, PWCNet optimizes the efficiency of dynamic layout, and MOFA strengthens the aesthetic and spatial presentation. Although the complete M-SVP model has the best comprehensive performance, its inference time is nearly double that of the baseline model, reflecting the computational cost of integrated multi-module intelligence and highlighting the trade-off between performance gain and computational overhead.

To verify the generalization ability of the M-SVP model in unseen digital media scenarios, two types of external datasets with significant differences from the training set scenarios were selected: the E-commerce Product Detail Page dataset [26], containing 50 000 samples, covering pages of clothing, electronics, and food, characterized by a dense arrangement of *multiple images + short text + price tags*, and the Social Media Post dynamic dataset [25], containing 30 000 samples, covering WeChat official accounts and Weibo images and text, characterized by *irregular layout + mixed emoticons/topic tags*). On the above datasets, the core metrics of M-SVP were compared with those of GA-YOLOv8, Mask R-CNN, and ResNet-50, and the results are shown in Table 4. These results indicate that M-SVP still maintains high performance in two types of unfamiliar scenarios: the extraction accuracy rate exceeds 89% in both cases, mIOU is ≥ 0.84 , the aesthetic score and space occupancy rate only decrease by 5–8% compared with the training set, while the performance degradation of the comparison models generally reaches 15–25%. In conclusion, M-SVP demonstrates strong generalization ability in cross-scenario tasks, verifying its practicality as a general digital media processing solution.

To visually present the focus of attention and correlation logic of the model layout decision, the typical digital media scenario of the social media page is still selected. The specific layout decision process of M-SVP is analyzed through the attention weight graph, and the result is shown in Figure 12. This Figure shows the attention weight

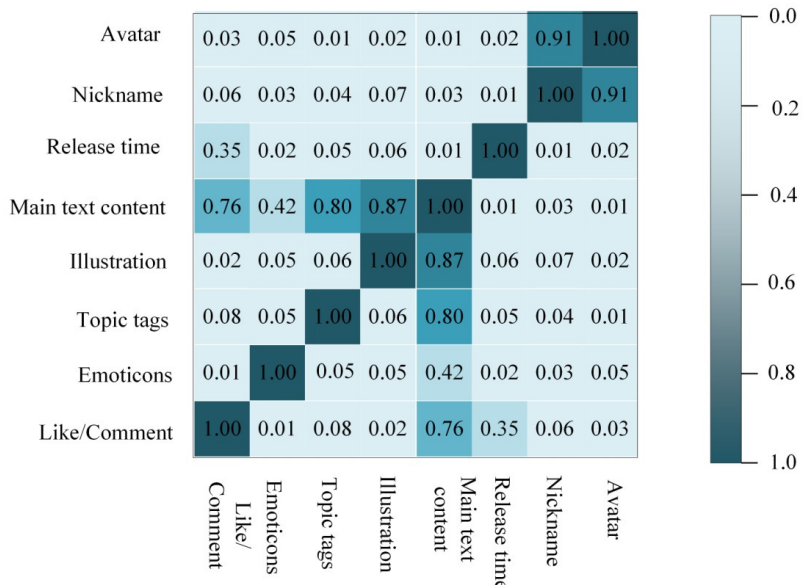


Fig. 13. Social media page model layout attention weight chart.

heat map of eight core elements in social media pages. The correlation strength between elements is quantified through color depth and numerical values. Among them, the correlation degree between avatars and nicknames exceeds 0.9, and a strong binding of *visual identity* + *text identity* is used to build a fast recognition channel for users about the identity of the publisher. The main text is deeply associated with the accompanying images (0.87) and topic tags (0.80), prioritizing the strengthening of the collaborative relationship between *text information* – *visual supplementation* – *dissemination classification* to meet the dissemination needs of *efficient content reach* on social media. The correlation between the main text and the interactive area (0.76) is particularly significant. Through the connection of *core content* → *interactive feedback*, the willingness to interact is strengthened, and it precisely alters the scene behavior pattern of *emotion-driven interaction*. The overall weight distribution clearly echoes the scene function chain of *identity recognition* – *content understanding* – *interactive dissemination*, intuitively verifying the scene pertinence and logical rationality of the model layout decision.

5. Conclusion

To address the issues of fuzzy visual element extraction and layout optimization, which rely on manual experience and lead to a balance problem between efficiency and aesthetics in digital media, the study designed the M-SVP model. This model constructs a multimodal collaborative architecture, covering core modules for precise visual element extraction, dynamic correlation analysis, and intelligent layout optimization, offering an intelligent solution for automated digital media design. The experimental results showed that the SAM-ViT algorithm achieved the highest visual element extraction accuracy of 98.8% across different categories. As the number of training iterations increased to 50, its mIoU value stabilized at 0.95, and its F1 score reached a maximum of $98.17 \pm 0.38\%$. Furthermore, the M-SVP model demonstrated 99% extraction accuracy for easily confused visual elements. After layout optimization, the M-SVP model's aesthetic score improved to 95.6, and its spatial occupancy rate increased to 97.2%, far exceeding the comparison models. In conclusion, the M-SVP model exhibited excellent performance in information extraction, layout optimization, and robustness testing. However, the model itself still has certain limitations. Its multi-module integration leads to high computational complexity and insufficient real-time performance when deployed on edge devices with limited computing power. Future work will focus on optimizing the above-mentioned deficiencies, compressing model parameters through knowledge distillation to reduce computational costs, and further enhancing the practicality and universality of the model.

Fundings

The research is supported by Shanxi Provincial Education Science *14th Five-Year Plan* 2022 annual general planning project *New Era of higher vocational colleges Computer Basics course ideological and political implementation path* (project number: GH-220519); Shanxi Engineering Vocational College 2022 annual teaching and research projects *Research on University Computer Basic Teaching Based on Computational Thinking* (Project number: JY2022-12).

References

- [1] J. D. Blair, K. M. Gaynor, M. S. Palmer, and K. E. Marshall. A gentle introduction to computer vision-based specimen classification in ecological datasets. *Journal of Animal Ecology* 93(2):147–158, 2024. doi:10.1111/1365-2656.14042.
- [2] F. Chen, L. Chen, H. Han, S. Zhang, D. Zhang, et al. The ability of Segmenting Anything Model (SAM) to segment ultrasound images. *BioScience Trends* 17(3):211–218, 2023. doi:10.5582/bst.2023.01128.

- [3] J. Deng, A. Berg, S. Satheesh, H. Su, A. Khosla, et al. ImageNet Large Scale Visual Recognition Challenge 2012 (ILSVRC2012). IMAGENET, 2012. <https://www.image-net.org/challenges/LSVRC/2012/>.
- [4] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, et al. ImageNet: A large-scale hierarchical image database. In: *2009 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 248–255, 2009. doi:10.1109/CVPR.2009.5206848.
- [5] J.-H. Ha and H. Lee. A deep learning model for precipitation nowcasting using multiple optical flow algorithms. *Weather and Forecasting* 39(1):41–53, 2024. doi:10.1175/WAF-D-23-0104.1.
- [6] E. Hassan, M. Y. Shams, N. A. Hikail, and S. Elmougy. The effect of choosing optimizer algorithms to improve computer vision tasks: A comparative study. *Multimedia Tools and Applications* 82(11):16591–16633, 2023. doi:10.1007/s11042-022-13820-0.
- [7] A. J. Jimeno Yepes. PubLayNet. GitHub, 2025. <https://github.com/ibm-aur-nlp/PubLayNet>.
- [8] S. Khoubani and M. H. Moradi. A deep learning phase-based solution in 2D echocardiography motion estimation. *Physical and Engineering Sciences in Medicine* 47(4):1691–1703, 2024. doi:10.1007/s13246-024-01481-2.
- [9] S. Kitada. huggingface-datasets_Magazine. GitHub, 2023. https://github.com/creative-graphic-design/huggingface-datasets_Magazine.
- [10] A. Krizhevsky. *Learning Multiple Layers of Features from Tiny Images*. Master’s thesis, University of Toronto, 2009. <https://www.cs.toronto.edu/~kriz/learning-features-2009-TR.pdf>.
- [11] A. Krizhevsky, V. Nair, and G. Hinton. The CIFAR-10 dataset. In Alex Krizhevsky home page, 2009. <https://www.cs.toronto.edu/~kriz/cifar.html>.
- [12] M. Y. Landolsi, L. Hlaoua, and L. Ben Romdhane. Information extraction from electronic medical documents: state of the art and future research directions. *Knowledge and Information Systems* 65(2):463–516, 2023. doi:10.1007/s10115-022-01779-1.
- [13] C. Li, Y. Huang, W. Li, H. Liu, X. Liu, et al. Flaws can be applause: Unleashing potential of segmenting ambiguous objects in SAM. *Advances in Neural Information Processing Systems* 37:45578–45599, 2024. doi:10.52202/079017-1449.
- [14] J. Li, Z. Zhou, J. Yang, A. Pepe, C. Gsaxner, et al. MedShapeNet – a large-scale dataset of 3D medical shapes for computer vision. *Biomedical Engineering / Biomedizinische Technik* 70(1):71–90, 2025. doi:10.1515/bmt-2024-0396.
- [15] H. B. Mahajan, N. Uke, P. Pise, M. Shahade, V. G. Dixit, et al. Automatic robot manoeuvres detection using computer vision and deep learning techniques: a perspective of internet of robotics things (IoRT). *Multimedia Tools and Applications* 82(15):23251–23276, 2023. doi:10.1007/s11042-022-14253-5.
- [16] T. Onyejelem and A. Eric Msughter. Digital generative multimedia tool theory (DGMTT): A theoretical postulation. *Journalism and Mass Communication* 14(3):189–204, 2024. doi:10.17265/2160-6579/2024.03.004.
- [17] A. S. Ortega-Calvo, R. Morcillo-Jimenez, C. Fernandez-Basso, K. Gutiérrez-Batista, M. A. Vila, et al. AIMDP: An artificial intelligence modern data platform. Use case for Spanish national health service data silo. *Future Generation Computer Systems* 143:248–264, 2023. doi:10.1016/j.future.2023.02.002.
- [18] E. W. Prastyaningtyas, A. M. A. Ausat, L. F. Muhamad, M. I. Wanof, and S. Suherlan. The role of information technology in improving human resources career development. *Jurnal Teknologi Dan Sistem Informasi Bisnis* 5(3):266–275, 2023. doi:10.47233/jteksis.v5i3.870.

- [19] B. Rokh, H. Mirvaziri, and M. H. Olyae. A new evolutionary optimization based on multi-objective firefly algorithm for mining numerical association rules. *Soft Computing* 28(9):6879–6892, 2024. doi:[10.1007/s00500-023-09558-y](https://doi.org/10.1007/s00500-023-09558-y).
- [20] O. Russakovsky, J. Deng, H. Su, J. Krause, S. Satheesh, et al. ImageNet Large Scale Visual Recognition Challenge. *International Journal of Computer Vision* 115(3):211–252, 2015. doi:[10.1007/s11263-015-0816-y](https://doi.org/10.1007/s11263-015-0816-y).
- [21] S. R. Shen, J. Balakrishnan, and C. H. Cheng. Dynamic content layout optimization for news website front pages. *Journal of Modelling in Management* 19(6):1907–1926, 2024. doi:[10.1108/JM2-01-2024-0015](https://doi.org/10.1108/JM2-01-2024-0015).
- [22] K. Subramanian, F. Hajamohideen, V. Viswan, N. Shaffi, and M. Mahmud. Exploring intervention techniques for Alzheimer’s disease: Conventional methods and the role of AI in advancing care. *Artificial Intelligence and Applications* 2(2):59–77, 2024. doi:[10.47852/bonview42022497](https://doi.org/10.47852/bonview42022497).
- [23] D. Sun, X. Yang, M.-Y. Liu, and J. Kautz. PWC-Net: CNNs for optical flow using pyramid, warping, and cost volume. In: *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 8934–8943, 2018. doi:[10.1109/CVPR.2018.00931](https://doi.org/10.1109/CVPR.2018.00931).
- [24] K. Tan, J. Wu, H. Zhou, Y. Wang, and J. Chen. Integrating advanced computer vision and AI algorithms for autonomous driving systems. *Journal of Theory and Practice of Engineering Science* 4(1):41–48, 2024. doi:[10.53469/jtpes.2024.04\(01\).06](https://doi.org/10.53469/jtpes.2024.04(01).06). <https://centuryscipub.com/index.php/jtpes/article/view/427>.
- [25] P. Tank. Social media post dataset. Kaggle Dataset, 2024. <https://www.kaggle.com/datasets/prishatank/post-generator-dataset/data>.
- [26] F. Tiago. ecommerce-product-dataset. GitHub Repository, 2025. <https://github.com/octaprice/ecommerce-product-dataset>.
- [27] D. Wang, J. Zhang, B. Du, M. Xu, L. Liu, et al. SAMRS: Scaling-up remote sensing segmentation dataset with segment anything model. *Advances in Neural Information Processing Systems* 36:8815–8827, 2023. doi:[10.48550/arXiv.2305.02034](https://doi.org/10.48550/arXiv.2305.02034).
- [28] Y. Xue and J. Williams. Inducing shifts in attentional and preattentive visual processing through brief training on novel grammatical morphemes: An event-related potential study. *Language Learning* 74(S1):185–223, 2024. doi:[10.1111/lang.12642](https://doi.org/10.1111/lang.12642).
- [29] P. Zhang, J. Zheng, H. Lin, C. Liu, Z. Zhao, et al. Vehicle trajectory data mining for artificial intelligence and real-time traffic information extraction. *IEEE Transactions on Intelligent Transportation Systems* 24(11):13088–13098, 2023. doi:[10.1109/TITS.2022.3178182](https://doi.org/10.1109/TITS.2022.3178182).
- [30] X. Zheng, X. Qiao, Y. Cao, and R. W. Lau. Content-aware generative modeling of graphic design layouts. *ACM Transactions on Graphics* 38(4):133, 2019. doi:[10.1145/3306346.3322971](https://doi.org/10.1145/3306346.3322971).
- [31] X. Zhong, J. Tang, and A. Jimeno Yepes. PubLayNet: Largest dataset ever for document layout analysis. In: *2019 International Conference on Document Analysis and Recognition (ICDAR)*, pp. 1015–1022, 2019. doi:[10.1109/ICDAR.2019.00166](https://doi.org/10.1109/ICDAR.2019.00166).

