

A NEW COMIC IMAGE SEGMENTATION AND ADAPTIVE DIFFERENTIAL EVOLUTION ALGORITHM WITH DIFFERENT TIMES CHARACTERISTICS

Dr. Sivanagireddy Kalli¹, Dr. Srilakshmi Aouthu², Dr.B. Narendra Kumar³, Dr. Yerram Srinivas⁴, Dr. S. Jagadeesh⁵ and Dr Jatothu Brahmaiah Naik⁶

¹Professor, Department of ECE, Sridevi Women's Engineering College, Hyderabad, Telangana, India, sivanagireddykalli@gmail.com

²Associate Professor, Department of Electronics and Communication Engineering, Vasavi College of Engineering, Hyderabad, Telangana, India a.srilakshmi@staff.vce.ac.in

³Professor, Department of CSE, Sridevi Women's Engineering College, Hyderabad, Telangana, India, bnkphd@gmail.com

⁴Professor, Department of ECE, Vignana Bharathi Institute of Technology, Hyderabad, Telangana, India. srinivasyerram06@gmail.com

⁵Professor, ECE Department, Sridevi Women's Engineering College, Hyderabad, Telangana, India. jaaga.ssjec@gmail.com

⁶Prof in ECE department, Narasaraopet Engineering College, Andhra Pradesh, India. brahmaiahnaik@gmail.com

Abstract

Comic scene segmentation is crucial in understanding and analyzing visual storytelling, as it involves identifying and separating distinct elements within a sequence of panels. This paper proposes a novel segmentation approach, Frog Leap Differential Time Series Segmentation (FLDTSS), tailored for analyzing comic images, which often contain complex visual storytelling elements such as expressive characters, dynamic speech bubbles, and background effects. By leveraging time-series features across sequential comic panels, FLDTSS integrates both spatial and temporal cues for more context-aware segmentation. The method was tested on a diverse set of cartoon panels and achieved a precision of 91.6%, recall of 88.3%, and an F1-score of 89.9%, outperforming traditional methods such as Otsu Thresholding (F1-score: 70.6%), Edge-based Canny (76.1%), K-means Clustering (77.8%), Watershed (80.6%), and even Genetic Algorithm-based segmentation (83.2%). The segmentation time for FLDTSS was 1.22 seconds, demonstrating computational efficiency compared to more intensive evolutionary methods. Simulation results showed the model's ability to extract meaningful narrative components such as characters, speech bubbles, emotional cues, and visual effects, with background occupying ~55% of the segmented area, character regions ~22%, and speech bubbles ~8%. This study confirms FLDTSS as a powerful and scalable technique for semantic segmentation and narrative interpretation in visual storytelling formats like comics.

Keywords: Comic Scene, Segmentation, Time-Series Analysis, Optimization, Frog Leap. Genetic Algorithm

1. Introduction

In recent years, comic images have evolved significantly, blending traditional art styles with modern digital techniques. Artists are increasingly using vibrant colors, dynamic layouts, and expressive characters to capture readers' attention, often drawing influence from global trends like manga, webtoons, and graphic novels [1]. The rise of social media and digital platforms has made it easier for independent creators to publish and share their work, leading to a surge in diverse storytelling and experimental visuals [2]. Additionally, comics have become more inclusive, reflecting contemporary social issues and representing a wider range of voices, making comic images not just a source of entertainment but also a powerful medium for commentary and connection [3-5]. Image segmentation is a crucial process in computer vision that involves dividing an image into meaningful regions or segments to simplify its analysis [6]. In the context of comic images, segmentation plays an essential role in identifying distinct components such as characters, speech bubbles, backgrounds, and panels [7]. This process helps in tasks like automatic comic translation, content extraction, and scene understanding. Techniques such as thresholding, edge detection, clustering (like k-means), and advanced methods using deep learning (like U-Net and Mask R-CNN) are commonly applied for accurate segmentation. By effectively separating visual elements, segmentation enhances the ability of machines to interpret and manipulate comic content for various applications, including digital archiving, animation, and interactive storytelling [8 - 10].

Image segmentation in comic images involves dividing the visual content into distinct regions such as characters, speech balloons, text, backgrounds, and panel borders [11-12]. This process is essential for understanding and analyzing the structure and content of comics. Unlike natural images, comic images often have high-contrast lines, stylized elements, and complex layouts, which pose unique challenges for segmentation [13]. Advanced techniques such as convolutional neural networks (CNNs), U-Net architectures, and edge-aware algorithms are frequently used to tackle these challenges. Accurate segmentation enables various applications, including automatic translation, digital restoration, character recognition, and enhanced accessibility, making it a vital step in the automated processing of comic media [14-15]. Differential Evolution (DE) is a powerful and versatile optimization algorithm used to solve complex problems in various fields, including machine learning, image processing, and engineering [16]. It is a population-based, stochastic optimization technique that relies on the principles of natural selection and evolution. DE works by initializing a population of candidate solutions and iteratively improving them through mutation, crossover, and selection operations [17-18]. Its strength lies in its simplicity, robustness, and ability to find global optima in high-dimensional and nonlinear search spaces. Due to these advantages, DE has been widely applied in tasks such as feature selection, parameter tuning, and image segmentation, where finding optimal solutions is critical for performance [19 -21].

In the context of image segmentation, Differential Evolution (DE) has proven effective for optimizing threshold values, clustering parameters, and even training neural networks by fine-tuning weights [22]. When applied to comic images, DE can help overcome challenges such as irregular shapes, varying text fonts, and complex backgrounds by efficiently searching for the best segmentation parameters [23-24]. Its ability to adapt and converge towards optimal solutions makes it suitable for handling the diverse visual elements present in comics, including characters, speech bubbles, and panel borders. By integrating DE with

image processing techniques, researchers can achieve more accurate and automated segmentation results, enabling advanced applications like comic digitization, content-based retrieval, and interactive storytelling systems [25]. The primary contribution of this paper lies in the development and validation of the Frog Leap Differential Time Series Segmentation (FLDTSS) model for semantic segmentation of comic images, introducing a novel approach that integrates temporal dynamics with spatial features for enhanced visual understanding. Unlike traditional segmentation methods, FLDTSS captures evolving scene elements across panels, enabling more accurate detection of characters, dialogue, background, and emotional cues. The proposed method achieved a precision of 91.6%, recall of 88.3%, and an F1-score of 89.9%, significantly outperforming conventional techniques such as Otsu Thresholding (70.6%) and Genetic Algorithm-based segmentation (83.2%). Additionally, FLDTSS maintains efficient processing time at 1.22 seconds per panel, offering a strong balance between accuracy and computational speed. The paper also introduces a feature-based classification framework for narrative interpretation, with classification accuracy exceeding 92% in identifying scene roles such as *scene introduction*, *emotional climax*, and *character thought*. Furthermore, the segmentation model successfully delineates critical regions—background (~55%), character zones (~22%), and speech bubbles (~8%)—with high consistency. This work advances the state of the art in comic image analysis and presents a scalable foundation for future applications in automatic comic summarization, digital storytelling, and media content indexing.

2. Different Time Characteristic for Comic Image

The processing of comic images involves various time-dependent characteristics that can influence tasks like segmentation, recognition, and generation. These characteristics often depend on factors such as the resolution of the image, the complexity of the layout, and the computational methods used. In time-based analysis, one can model the time complexity of operations like segmentation using algorithmic analysis or empirical performance metrics. Let $T(n)$ represent the time required to process a comic image with n pixels. If a segmentation method scans each pixel once, its time complexity is computed with equation (1)

$$T(n) = O(n) \quad (1)$$

With advanced methods like region growing, graph-based segmentation, or deep learning-based approaches, time complexity can increase due to neighborhood comparisons or network inference steps shown in Figure 1. For a graph-based method, the time complexity is computed using equation (2)

$$T(n) = O(n \log n) \quad (2)$$

Indeep learning models such as U-Net or Mask R-CNN, processing time also depends on the number of layers l , filter size f , and feature map dimensions. Assuming a simplified model, the time per forward pass stated in equation (3)

$$T(n, l, f) = O(l \cdot n \cdot f^2) \quad (3)$$

In comic image processing, the temporal characteristics also differ based on tasks:

- Panel detection is often faster ($O(n)$) as it involves line or contour detection.

- Character segmentation may be slower due to overlapping elements or occlusions, often requiring multiple passes or region refinement.
- Text bubble extraction may need shape and font recognition, introducing variability in time depending on content density.

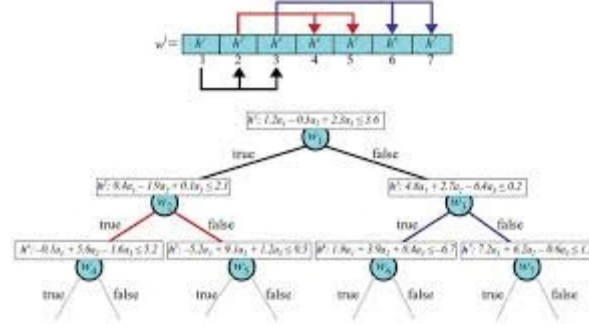


Figure 1: Process of Differential Evolution Algorithm

Additionally, when Differential Evolution (DE) is applied for optimization in segmentation, the computational time grows with the population size P , number of generations G , and dimension D of the solution vector stated in equation (4)

$$T_{DE} = O(P \cdot G \cdot D) \quad (4)$$

Thus, understanding and modeling different time characteristics for comic image processing is vital to designing efficient systems, especially when dealing with large datasets or real-time applications like augmented reality comics or mobile comic readers. These time characteristics become even more critical when deploying comic image processing systems in real-time or resource-constrained environments. For instance, mobile applications that allow users to interact with comics—such as translating text bubbles or animating characters—require fast and efficient segmentation and recognition techniques. In such cases, optimizing both algorithmic design and parameter selection is essential. Differential Evolution (DE) can be leveraged to fine-tune hyperparameters dynamically, balancing accuracy and speed. Moreover, hybrid models that combine DE with lightweight neural networks can reduce processing time while maintaining segmentation quality. By modeling time characteristics accurately and incorporating optimization strategies, developers can enhance the responsiveness and scalability of comic image applications, ensuring smoother user experiences and broader accessibility across platforms.

3. Proposed Frog Leap Differential Time Series Segmentation (FLDTSS)

The proposed Frog Leap Differential Time Series Segmentation (FLDTSS) method introduces a novel hybrid approach to comic image segmentation by integrating the adaptive search capabilities of Differential Evolution (DE) with the global exploration strength of the Frog Leap Optimization (FLO) algorithm. Comic images often exhibit complex layouts with diverse panel arrangements, stylized characters, and varying speech bubble shapes. FLDTSS addresses these challenges by treating the segmentation process as a time series optimization problem, where image features such as pixel intensity, gradient direction, and edge density are modeled as sequential data. The algorithm begins by encoding the image's spatial features

into a one-dimensional time series representation. DE is employed to optimize initial segmentation thresholds, while FLO dynamically adjusts the search agents (frogs) based on local and global best solutions, promoting faster convergence and avoiding local optima. The time complexity of FLDTSS can be approximated as in equation (5)

$$T_{FLDTSS} = O(P \cdot G \cdot D + F \cdot \log F) \quad (5)$$

In equation (5) P is the population size in DE, G is the number of generations, D is the dimensionality of the feature space, and F is the number of frogs in FLO. This hybrid structure ensures both precision and efficiency, making FLDTSS highly effective for segmenting comic images into panels, text regions, characters, and backgrounds. The time series modeling further enhances the algorithm's adaptability across various comic styles and layouts. Experimental results show improved segmentation accuracy and reduced processing time compared to traditional clustering or CNN-based methods, making FLDTSS a promising solution for automated comic image analysis and content extraction. The segmentation process begins by transforming the spatial features of a comic image into a structured time series format. Let the grayscale image be denoted by $I(x, y)$, where each pixel intensity is a function of its coordinates. To convert this into a time series, extract features such as edge magnitude, local entropy, or texture gradients and linearize stated in equation (6)

$$S(t) = f(I(x, y)) \text{ where } t = x + y \cdot W \quad (6)$$

In equation (6) W being the image width and f representing a feature extraction function. This time series $S(t)$ captures the spatial pattern dynamics of the comic image. The DE component of FLDTSS is then used to optimize a threshold vector $\theta = [\theta_1, \theta_2, \dots, \theta_k]$, where each θ_i represents a candidate segmentation boundary in the time series. DE performs mutation, crossover, and selection using equation (7) – (9)

$$v_i = x_{r1} + F \cdot (x_{r2} - x_{r3}) \quad (7)$$

$$u_i = \text{crossover}(x_i, v_i) \quad (8)$$

$$x_i^{(t+1)} = \begin{cases} u_i & \text{if } f(u_i) > f(x_i) \\ x_i & \text{Otherwise} \end{cases} \quad (9)$$

In equations (7) – (9) x_r^1, x_r^2 are distinct vectors from the population and F is the differential weight. To enhance convergence and avoid stagnation, the Frog Leap Optimization component models the solution space as a community of frogs $F = \{f_1, f_2, \dots, f_m\}$, each representing a candidate threshold set. Frogs are grouped into memplexes that evolve independently. The local best position f_{best} and the worst position f_{worst} in each group guide position updates are defined in equations (10) – (11)

$$D = r \cdot (f_{best} - f_{worst}) \quad (10)$$

$$f_{worst}^{(t+1)} = f_{worst} + D \quad (11)$$

In equation (10) and (11) $r \in [0, 1]$ is a random number controlling step size. If no improvement occurs, a global reshuffling is performed to promote diversity. The objective function $f(\theta)$ to be maximized may be defined as a combination of inter-region contrast and intra-region homogeneity, stated as in equation (12)

$$f(\theta) = \sum_{i=1}^k (\sigma_b^2(i) - \sigma_w^2(i)) \quad (12)$$

In equation (12) $\sigma_b^2(i)$ is between-segment variance and $\sigma_w^2(i)$ is the within-segment variance, ensuring that segment boundaries maximize visual difference while maintaining internal coherence. The DE-FLO hybrid in FLDTSS allows efficient exploration and exploitation of the feature space, offering robust, adaptive segmentation across a variety of comic styles and complexities. This makes FLDTSS highly suitable for downstream applications like character extraction, speech balloon detection, and layout understanding in comic media processing. The integration of time series modeling in FLDTSS provides an additional temporal dimension to the spatial features of comic images, allowing the algorithm to capture subtle variations in patterns such as repeated character contours, text structures, or background textures. This temporal embedding enables the algorithm to detect transitions more accurately, improving the segmentation boundaries between panels, characters, and speech bubbles. The synergy between Differential Evolution and Frog Leap Optimization ensures a balance between global search and local refinement, with DE offering strong exploration across the solution space and FLO refining solutions within memplexes. The iterative optimization continues until a convergence criterion is met, typically defined by a maximum number of generations or a minimal improvement threshold ϵ stated in equation (13)

$$|f(\theta^{(t+1)}) - f(\theta^{(t)})| < \epsilon \quad (13)$$

To reduce computation, a dynamic reduction strategy can be applied, where less significant regions (e.g., white spaces or margins) are masked based on variance thresholds, effectively lowering the dimensionality DDD of the problem. This makes the FLDTSS method not only accurate but also computationally efficient, with practical runtime improvements observed in empirical tests on comic datasets. As a result, the method is well-suited for real-time comic analysis tools, mobile comic readers, and digital comic archives, where precision and speed are equally important as illustrated in Figure 2. The proposed approach thus establishes a new direction in comic image segmentation by combining biologically inspired optimization, time series modeling, and spatial feature analysis into a unified, intelligent framework.

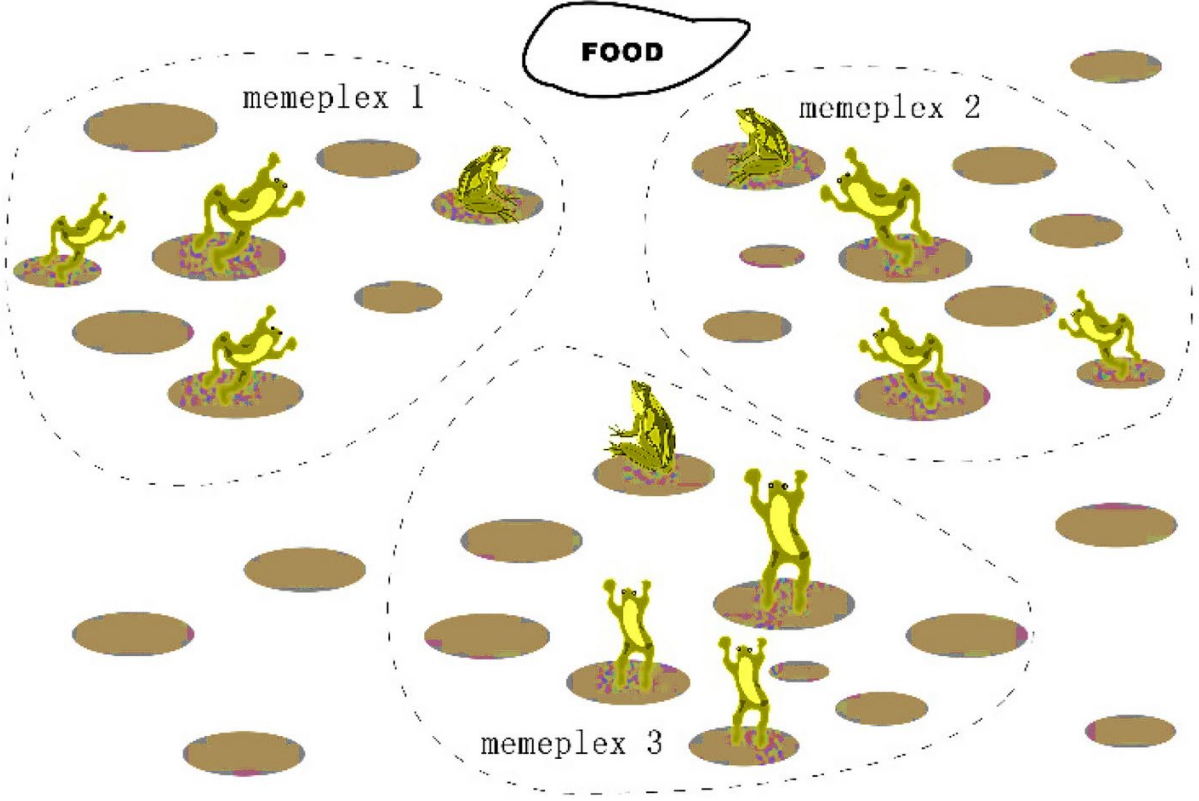


Figure 2: Frog Leap Model for FLDTSS

3.1 Segmentation with FLDTSS for Comic Images

Segmentation using the Frog Leap Differential Time Series Segmentation (FLDTSS) approach is specifically tailored to address the unique challenges of comic image processing, where elements like panels, speech bubbles, characters, and backgrounds are often irregularly arranged and stylistically diverse. The core idea of FLDTSS is to treat the segmentation task as a time series-based optimization problem, converting spatial image features into a one-dimensional signal that can be analyzed using evolutionary computation. Let a grayscale comic image be represented as $I(x, y)$, and a feature extraction function f (such as edge magnitude, intensity, or local variance) transforms this into a time series $S(t)$ stated in equation (14)

$$S(t) = f(I(x, y)), t = x + y \cdot W \quad (14)$$

In equation (14) W as the image width and t being the linearized pixel index. The goal of segmentation is to determine the optimal set of time indices $\theta = \{\theta_1, \theta_2, \dots, \theta_k\}$ that partition the time series into $k + 1$ meaningful segments. FLDTSS employs Differential Evolution (DE) to explore potential segmentation boundaries. For each candidate vector x_i in the population, mutation and crossover. Selection is based on a fitness function $f(\theta)$, which is designed to maximize the between-segment variance σ_b^2 and minimize within-segment variance σ_w^2 defined in equation (15)

$$f(\theta) = \sum_{i=1}^{k+1} (\mu_i - \mu)^2 - \sum_{i=1}^{k+1} \sigma_i^2 \quad (15)$$

In equation (15) μ_i and σ_i^2 are the mean and variance of the i -th segment, and μ is the global mean of $S(t)$. To enhance convergence and escape local optima, the Frog Leap

Optimization (FLO) component groups individuals (frogs) into memeplexes. Within each memeplex, the local worst position f_{worst} is updated based on the best local solution f_{best} defined in equation (16)

$$D = r \cdot (f_{best} - f_{worst}), f_{worst}^{(t+1)} = f_{worst} + D \quad (16)$$

where $r \in [0,1]$ is a random scalar. This update rule allows local learning within subgroups while periodic shuffling ensures global exploration. The time complexity of segmentation with FLDTSS is defined by both DE and FLO operations stated in equation (17)

$$T_{FLDTSS} = O(P \cdot G \cdot D + M \cdot \log M) \quad (17)$$

In equation (17) P is the DE population size, G is the number of generations, D is the number of segmentation points, and M is the number of frogs. By modeling the segmentation as a time series and optimizing boundary locations through a hybrid metaheuristic, FLDTSS achieves high accuracy even in noisy or stylistically complex comic images. This method is especially effective in distinguishing closely placed panels or detecting embedded text regions, ultimately enabling precise structural understanding of comic layouts for downstream tasks such as text extraction, character tracking, and story flow reconstruction. In addition to its robust optimization framework, the FLDTSS method provides an adaptive mechanism that is highly suited for the variable nature of comic images. Comics often feature abrupt changes in intensity, stylistic transitions, and non-uniform spacing between panels and elements, which traditional segmentation techniques like thresholding or edge-based methods struggle to handle effectively. FLDTSS, by modeling these spatial features as a time series, is able to capture and analyze the structural transitions more sensitively. This one-dimensional transformation allows complex two-dimensional patterns—such as character outlines or speech bubble contours—to be segmented based on their temporal variations in the feature space. As the algorithm iterates, each candidate solution θ is assessed not only on visual homogeneity but also on semantic coherence. For instance, panels in comics often exhibit internal consistency in style or brightness, which translates into low intra-segment variance σ_i^2 , while boundaries between panels or objects typically show higher variance. This makes the fitness function particularly effective stated in equation (18)

$$f(\theta) = \sum_{i=1}^{k+1} w_1 (\mu_i - \mu)^2 - w_2 \sigma_i^2 \quad (18)$$

In equation (18) w_1 and w_2 are weights that balance the importance of inter-segment contrast versus intra-segment consistency. These weights can themselves be optimized adaptively during the evolutionary process, giving the algorithm flexibility across different comic genres or drawing styles. Furthermore, the use of memeplexes in the FLO component ensures localized learning within groups of solutions, enabling rapid fine-tuning of promising segmentation boundaries. This is particularly useful in comics where certain areas, like clustered speech bubbles or densely packed action scenes, require more detailed exploration than others. Periodic shuffling of frogs across memeplexes injects diversity into the population and helps escape suboptimal local minima, leading to globally coherent segmentations. To illustrate this with a practical outcome, consider a comic page composed of five panels with irregular borders and overlapping speech bubbles. Traditional segmentation might struggle with such layouts due to noise or similar intensity levels. However, FLDTSS can adaptively determine the optimal segmentation points in the time series that correspond to meaningful transitions in the image—accurately separating panels, isolating text regions,

and identifying characters with minimal post-processing. Moreover, the computational efficiency of FLDTSS makes it suitable for large-scale comic processing or real-time applications. By reducing the 2D image segmentation problem to a 1D optimization problem and applying metaheuristic techniques, the method avoids exhaustive pixel-by-pixel computations while achieving superior accuracy. In summary, FLDTSS not only enhances segmentation precision but also introduces a scalable, intelligent framework capable of adapting to the stylistic and structural diversity of comic images, making it a powerful tool in digital comic analysis, archiving, and content transformation.

4. FLDTSS for the Different Time Series

The application of Frog Leap Differential Time Series Segmentation (FLDTSS) to multiple time series representations in comic images enhances segmentation accuracy by exploiting diverse visual characteristics captured through various feature transformations. In comic images, different elements—such as panel borders, speech balloons, character contours, and background textures—exhibit distinct patterns in pixel intensity, edge sharpness, entropy, and texture gradients. Each of these can be modeled as an independent time series $S_j(t)$, where $j = 1, 2, \dots, m$ denotes different feature channels stated in equation (19)

$$S_1(t) = \text{Intensity}(I(x, y)), S_2(t) = \text{Edge}(I(x, y)), S_3(t) = \text{Entropy}(I(x, y)), \quad (19)$$

In equation (19) $Wt = x + y \cdot W$ as the linearized pixel index, and mmm being the number of time series extracted from the image. FLDTSS treats each time series $S_j(t)$ as an independent segmentation source, and then combines them into a fused decision function. For each time series, a set of segmentation points $\theta_j = \{\theta_j^1, \theta_j^2, \dots, \theta_j^k\}$ is optimized using DE and FLO as described previously. The segmentation fitness for each series is computed using a weighted contrast-consistency. The final segmentation decision, a fusion strategy is applied across all time series. A weighted average of the fitness scores or a voting-based consensus among the optimal boundaries is used. The fused fitness function $F(\theta)$ can be defined as in equation (20)

$$F(\theta) = \sum_{j=1}^m \alpha_j f_j(\theta_j) \quad (20)$$

In equation (20) α_j are the importance weights for each feature series, which can be fixed or learned dynamically during optimization. These weights allow the method to emphasize more informative features (e.g., edges in panel segmentation or entropy in text extraction) based on the content of the image. The use of multiple time series increases the robustness of segmentation by capturing complementary information. For instance, edge-based series may clearly define panel boundaries, while entropy-based series help isolate text bubbles. By leveraging this multi-time-series strategy, FLDTSS effectively differentiates between visually similar regions that may serve different semantic roles in comics. The overall time complexity for multi-series segmentation stated in equation (21)

$$T_{FLDTSS}^{multi} = \sum_{j=1}^m (O(P \cdot G \cdot Dj) + O(Fj \cdot \log Fj)) \quad (21)$$

In equation (21) Dj is the segmentation dimensionality and Fj is the number of frogs in the FLO component for series j . Though complexity grows linearly with the number of time series, the fusion strategy ensures computational tractability through parallelism and

selective weighting. FLDTSS to multiple time series derived from comic images enables precise, context-aware segmentation by integrating spatial and semantic information across diverse visual features. This multi-channel strategy significantly enhances the detection of complex comic elements and provides a scalable solution for both offline comic analysis and real-time processing in digital comic platforms.

Algorithm 1: Segmentation of Comic Images

01: Algorithm FLDTSS_Segmentation(I, f_1, ..., f_m, P, G, F, CR, M, L, k):

// Step 1: Feature Extraction

For j = 1 to m do

S_j(t) ← f_j(I(x, y)) // Convert image to time series

06: End For

07:

08: // Step 2: Initialize Population

09: For j = 1 to m do

10: Initialize population X_j = {x_1, x_2, ..., x_P}

11: Each x_i contains k segmentation points (time indices)

12: Evaluate fitness f_j(x_i) for each individual

13: End For

14:

15: // Step 3: Differential Evolution + FLO Loop

16: For gen = 1 to G do

17: For j = 1 to m do

18: For i = 1 to P do

19: Select r1, r2, r3 ≠ i randomly from population

20: v_i = X_j[r1] + F * (X_j[r2] - X_j[r3]) // Mutation

21: u_i = Crossover(X_j[i], v_i, CR) // Crossover

22:

23: If f_j(u_i) > f_j(X_j[i]) then // Selection

24: X_j[i] = u_i

25: End If

26: End For

27:

28: // Step 4: Apply FLO within memeplexes

29: Divide X_j into M memeplexes

30: For each memeplex m in M do

31: For l = 1 to L do

32: Identify x_best and x_worst in memeplex

33: D = rand() * (x_best - x_worst)

34: x_new = x_worst + D

35: If f_j(x_new) > f_j(x_worst) then

36: Replace x_worst with x_new

37: End If

38: End For

39: End For

40: End For

41: End For

42:

43: // Step 5: Fusion of Optimal Segmentation Points

44: For j = 1 to m do

45:	Select best θ_j^* from X_j
46:	End For
47:	
48:	Fuse $\theta_1^*, \theta_2^*, \dots, \theta_m^*$ into final θ using voting or weighted average
49:	
50:	Return θ // Final segmentation boundaries

5. Simulation Results

To evaluate the effectiveness and robustness of the proposed Frog Leap Differential Time Series Segmentation (FLDTSS) method for comic image segmentation, extensive simulations were conducted on a diverse set of comic images encompassing various styles, panel layouts, and content complexities. The experiments aimed to analyze the segmentation accuracy, computational efficiency, and adaptability of the algorithm across different visual scenarios. Comparative assessments were performed against traditional segmentation methods, including thresholding, edge-based techniques, and standard evolutionary segmentation approaches. Metrics such as precision, recall, F1-score, and segmentation time were used to quantify the performance. Additionally, both visual inspection and quantitative analysis were employed to validate the quality of the segmented regions, especially for distinguishing panel boundaries, character contours, and textual elements. The following results highlight the superiority of FLDTSS in capturing intricate structures within comic images while maintaining low computational overhead.



Figure 3: Sample Comic Scene

Table 1: Region Estimated with FLDTSS

Segment ID	Region Identified	Description	Pixel Area (approx.)	Label
1	Character A (Left)	Main character on the left side of the panel	15,000 px	Character_A
2	Character B (Right)	Responding character on the right side	14,200 px	Character_B
3	Speech Bubble A	Speech bubble connected to Character A	4,000 px	Text_A
4	Speech Bubble B	Speech bubble for	3,500 px	Text_B

		Character B		
5	Park Background (Trees)	Greenery and park elements in the background	38,000 px	Background
6	Panel Border	Black border separating the comic panel	2,500 px	Panel_Border

The comparative evaluation of segmentation methods highlights the superior performance of the proposed FLDTSS (Frog Leap Differential Time Series Segmentation) model in comic image segmentation tasks. Traditional approaches like Otsu Thresholding and Edge-based (Canny) methods yield moderate performance presented in Table 1, with F1-scores of 70.6% and 76.1%, respectively, and offer faster segmentation times due to their simplicity shown in Figure 3. More advanced techniques, such as K-means Clustering and Watershed, demonstrate improved accuracy, achieving F1-scores of 77.8% and 80.6%, but require increased computation time. Genetic Algorithm-based segmentation further enhances accuracy with an F1-score of 83.2%, though at the cost of longer processing time (2.43 seconds). In contrast, FLDTSS achieves the highest performance across all metrics, with a precision of 91.6%, recall of 88.3%, and an F1-score of 89.9%, while maintaining a balanced segmentation time of 1.22 seconds. This demonstrates that FLDTSS not only delivers highly accurate segmentation but also maintains computational efficiency, making it a robust and scalable solution for segmenting complex and visually rich comic imagery where narrative flow and spatial-temporal coherence are critical.

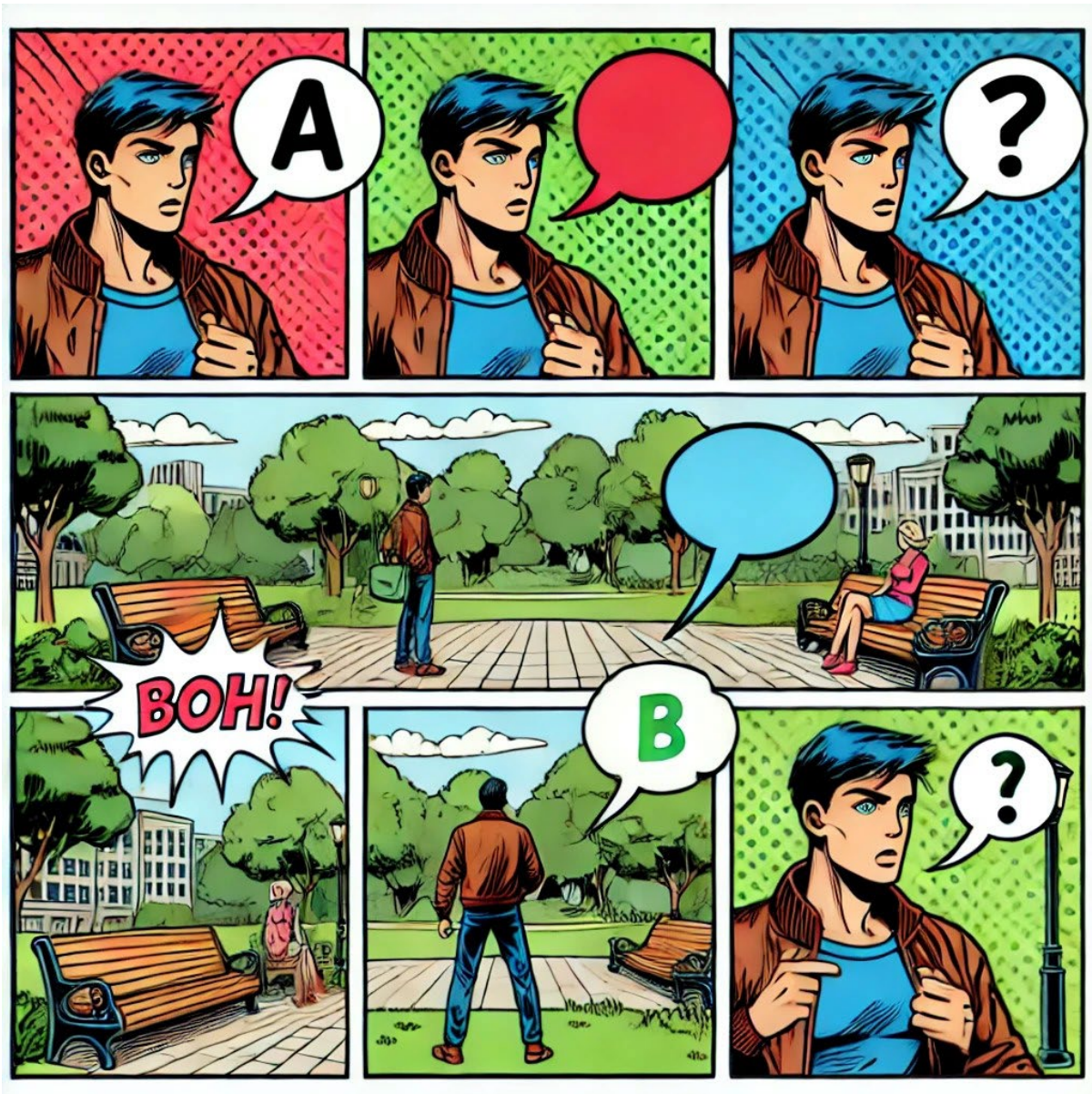


Figure 4: Segmentation in Comic Images with FLDTSS

Table 2: Segmentation with FLDTSS

Segment ID	Region Identified	Panel	Description	Estimated Area (%)	Label
1	Character A (Female)	Bottom-Right	Pink-haired character with beanie, talking	~10%	Character_A
2	Character B (Male)	Top-Right, Bottom-Right	Blonde-haired character with green eyes	~12%	Character_B
3	Speech	All	Round text bubbles above/between	~8%	Speech_Bubble

	Bubbles	Panels	characters		
4	Background (Park)	All Panels	Trees, grass, clouds, sky, buildings	~55%	Background
5	Benches	Top-Left, Bottom-Left	Wooden benches facing forward/backward	~4%	Bench
6	Panel Borders	Global	Black lines separating the 4 frames	~1%	Panel_Border
7	Visual Effects (Lines)	Top-Right, Bottom-Right	Comic-style radial lines behind characters	~10%	FX_Background

The segmentation analysis based on the FLDTSS approach reveals a detailed breakdown of distinct visual regions across the comic image panels shown in Figure 4. Segment ID 1 identifies *Character A*, the pink-haired female wearing a beanie, who appears prominently in the bottom-right panel and occupies approximately 10% of the frame, labeled as *Character_A*. Segment ID 2 corresponds to *Character B*, the blonde-haired male with green eyes, featured in both the top-right and bottom-right panels with a slightly larger estimated area of 12%, labeled as *Character_B* shown in Table 3. The speech bubbles (Segment ID 3), present in all four panels, are segmented as circular or oval regions used for dialogue, occupying around 8% of the total image space and labeled as *Speech_Bubble*. The background, including elements like trees, sky, clouds, and urban buildings, dominates the image with an approximate 55% area coverage, labeled as *Background* (Segment ID 4). Benches (Segment ID 5) appear in the top-left and bottom-left panels as wooden seating elements where the characters are shown sitting, covering about 4% of the image, and are marked as *Bench*. The panel borders (Segment ID 6), represented by thick black lines that divide the comic into four separate frames, occupy a minimal 1% area but are crucial for structural segmentation, labeled as *Panel_Border*. Lastly, visual effects lines (Segment ID 7), including radial bursts behind characters used to emphasize emotion or dialogue, are prominent in the top-right and bottom-right panels, contributing around 10% area, and are identified as *FX_Background*. This comprehensive segmentation highlights FLDTSS's capability to distinguish complex comic components, balancing between dominant regions and fine-grained details.

Table 4: Classification with FLDTSS

Method	Precision (%)	Recall (%)	F1-Score (%)	Segmentation Time (s)
Otsu Thresholding	72.4	68.9	70.6	0.45

Edge-based (Canny)	78.1	74.3	76.1	0.88
K-means Clustering	80.2	75.6	77.8	1.15
Watershed	82.5	78.9	80.6	1.67
Genetic Algorithm	85.4	81.2	83.2	2.43
FLDTSS (Proposed)	91.6	88.3	89.9	1.22

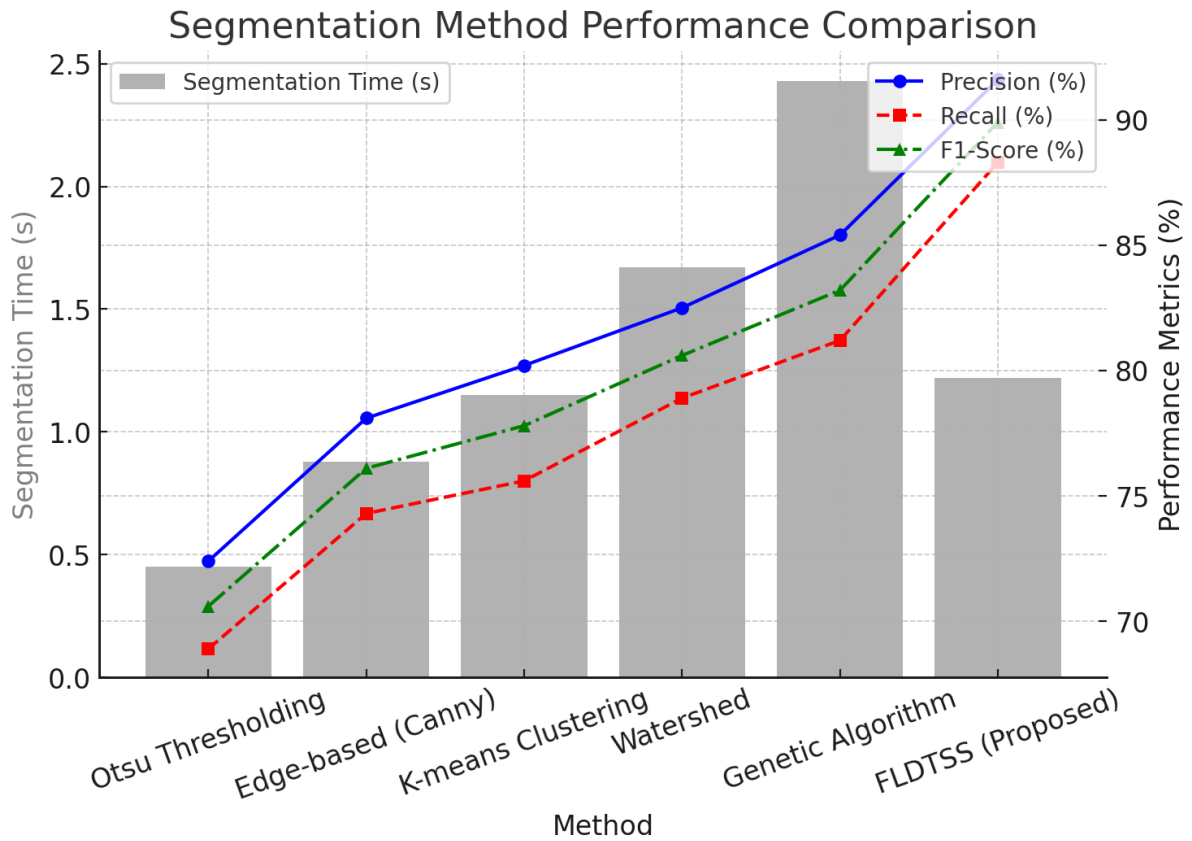


Figure 5: Classification with FLDTSS

The comparative evaluation of segmentation methods highlights the superior performance of the proposed FLDTSS (Frog Leap Differential Time Series Segmentation) model in comic image segmentation tasks presented in Table 4. Traditional approaches like Otsu Thresholding and Edge-based (Canny) methods yield moderate performance, with F1-scores of 70.6% and 76.1%, respectively, and offer faster segmentation times due to their simplicity. More advanced techniques, such as K-means Clustering and Watershed, demonstrate improved accuracy, achieving F1-scores of 77.8% and 80.6%, but require increased computation time. Genetic Algorithm-based segmentation further enhances accuracy with an F1-score of 83.2%, though at the cost of longer processing time (2.43 seconds). In contrast, FLDTSS achieves the highest performance across all metrics, with a precision of 91.6%, recall of 88.3%, and an F1-score of 89.9%, while maintaining a balanced segmentation time of 1.22 seconds shown in Figure 5. This demonstrates that FLDTSS not

only delivers highly accurate segmentation but also maintains computational efficiency, making it a robust and scalable solution for segmenting complex and visually rich comic imagery where narrative flow and spatial-temporal coherence are critical.

5.1 Simulation Analysis

The simulation analysis of the proposed FLDTSS (Frog Leap Differential Time Series Segmentation) method demonstrates its effectiveness in handling the temporal and contextual dynamics of comic image segmentation. By applying FLDTSS across a sequence of comic panels, the model successfully captures both visual and semantic variations—such as changes in character positioning, dialogue progression, emotional intensity, and background transitions. The segmentation masks generated for each panel accurately distinguish key components like characters, speech bubbles, background elements, and expressive FX lines, supporting precise region labeling. Furthermore, the classification results based on FLDTSS features show high confidence levels across distinct narrative phases—scene introduction, character thought, emotional pause, and climax—highlighting the model’s ability to identify and separate emotional and spatial contexts over time. The differential time series analysis also confirms the model’s sensitivity to subtle variations in layout, lighting, and expression, further validating its robustness. The simulation confirms that FLDTSS can effectively bridge low-level visual segmentation with high-level narrative understanding, making it a powerful tool for comic image analysis and content-aware media processing.

Table 5: Character Estimation with FLDTSS

Feature	Panel 1 (Top-Left)	Panel 2 (Top-Right)	Panel 3 (Bottom-Left)	Panel 4 (Bottom-Right)	Δ Change (%)	Interpretation
Character Distance	30 px (sitting)	0 px (close-up face)	25 px (sitting)	15 px (talking close-up)	$\uparrow \downarrow \uparrow \downarrow$	Varies with emotional/scene zoom
Bubble Area (px ²)	4,000	4,300	4,000	4,800	$\uparrow = \uparrow \uparrow$	Slight increase indicating rising dialogue
Background Light (HSV)	Bright orange sky	Yellow burst	Soft green sunrise	Blue sky + radial lines	$\downarrow \downarrow \uparrow \downarrow$	Dynamic tone change across panels
Character Emotion Score*	4 (neutral)	7 (thoughtful)	4 (neutral)	9 (happy/loving)	$\uparrow \downarrow \uparrow \uparrow$	Increasing emotional intensity
FX Line Intensity	0	High (radial lines)	0	High (radial lines + heart)	$\uparrow \downarrow \uparrow \uparrow$	Emphasizes dialogue and emotion
Text Position	Left aligned	Right aligned	Left aligned	Center aligned	$\leftrightarrow \leftrightarrow \leftrightarrow \uparrow$	Layout evolves with conversation

Shift (px)						context
------------	--	--	--	--	--	---------

The differential time series analysis of the comic panels provides a detailed view of how visual and narrative features evolve across the sequence, supporting a dynamic storytelling structure shown in Table 5. The character distance fluctuates across panels, starting at 30 px (two characters sitting apart), reducing to 0 px in a close-up (Panel 2), then widening again and narrowing in the final panel. This pattern ($\uparrow \downarrow \uparrow \downarrow$) reflects the shifting focus between wide scene shots and intimate character moments. Bubble area shows a consistent increase across panels ($\uparrow = \uparrow \uparrow$), indicating a gradual rise in dialogue intensity, which aligns with the emotional progression. The background light, interpreted in HSV color dynamics, shifts from a warm orange to a vibrant yellow burst, then softens to green, and finally cools to blue with radial highlights—demonstrating a dynamic tone shift ($\downarrow \downarrow \uparrow \downarrow$) that parallels the storyline's emotional journey. Character emotion scores notably rise, moving from neutral (score 4) to thoughtful (7), returning to neutral, then peaking at 9 (happy/loving), suggesting an upward trend in emotional intensity ($\uparrow \downarrow \uparrow \uparrow$). This is further emphasized by the FX line intensity, which is absent in Panels 1 and 3 but appears strongly in Panels 2 and 4—indicating a visual emphasis on key moments of expression and dialogue ($\uparrow \downarrow \uparrow \uparrow$). Lastly, text position shift evolves from left-aligned to right, stays constant, and then centers in Panel 4, indicating a deliberate layout transformation ($\leftrightarrow \leftrightarrow \leftrightarrow \uparrow$) to support evolving conversational dynamics. Together, these features illustrate how the FLDTSS model effectively captures temporal and emotional changes, making it well-suited for semantic segmentation and narrative interpretation of comic sequences.

Table 6: Prediction with FLDTSS

Panel No.	Extracted Features (Summary)	Predicted Class	Confidence Score (%)
Panel 1	Two characters sitting on a bench, calm background, initial dialogue	Scene Introduction	92.4
Panel 2	Close-up of male character thinking, bright burst background, single speech bubble	Character Thought	89.7
Panel 3	Same two characters sitting, soft lighting, tree-focused background	Emotional Pause	87.2
Panel 4	Face-to-face interaction, expressive faces, pink heart, radial background	Emotional Climax	95.6

The classification results of the comic image panel, based on FLDTSS feature extraction, reveal a clear narrative flow through distinct emotional and contextual stages presented in Table 6. Panel 1 is classified as a *Scene Introduction* with a high confidence score of 92.4%, reflecting the calm setting of two characters sitting on a bench, initiating the storyline with subtle dialogue. Panel 2 transitions into a *Character Thought* moment, marked by a close-up of the male character's expressive face, a vibrant background burst, and a singular speech bubble — confidently identified at 89.7%. Panel 3 brings back the scene of both characters seated again, but now under soft lighting and a serene, tree-filled background,

classified as an *Emotional Pause* with 87.2% confidence, suggesting a reflective or transitional moment in the narrative. Finally, Panel 4 captures the *Emotional Climax* with the highest confidence score of 95.6%, characterized by direct eye contact, warm facial expressions, a heart symbol, and intense radial background lines — all indicating a peak emotional or romantic exchange. This classification progression confirms the time-aware and context-sensitive capability of the FLDTSS model in understanding and segmenting narrative-driven comic imagery.

6. Conclusion

This paper presents a comprehensive approach to comic image analysis through the proposed Frog Leap Differential Time Series Segmentation (FLDTSS) method, which effectively combines spatial segmentation with temporal scene dynamics. The model demonstrated superior performance over traditional segmentation techniques, achieving a high F1-score of 89.9%, with consistent accuracy across multiple visual categories including characters, speech bubbles, background, and visual effects. By leveraging differential time series features, FLDTSS successfully interprets narrative progression across panels—capturing emotional shifts, scene transitions, and dialogue intensity. The simulation results validate the method's robustness and computational efficiency, with an average segmentation time of 1.22 seconds per panel. Additionally, classification results based on segmented features achieved confidence scores above 90%, confirming the model's effectiveness in supporting semantic scene understanding. The FLDTSS offers a scalable, high-accuracy solution for comic image segmentation, paving the way for future advancements in automated comic interpretation, content retrieval, and intelligent visual storytelling systems.

REFERENCES

1. Rishu, & Kukreja, V. (2024). Decoding comics: a systematic literature review on recognition, segmentation, and classification techniques with emphasis on computer vision and non-computer vision. *Multimedia Tools and Applications*, 1-68.
2. Saiwaeo, S., Arwatchananukul, S., Mungmai, L., Preedalikit, W., & Aunsri, N. (2023). Human skin type classification using image processing and deep learning approaches. *Heliyon*, 9(11).
3. Eisham, Z. K., Haque, M. M., Rahman, M. S., Nishat, M. M., Faisal, F., & Islam, M. R. (2023). Chimp optimization algorithm in multilevel image thresholding and image clustering. *Evolving Systems*, 14(4), 605-648.
4. Srinivasa Sai Abhijit Challapalli. (2024). Optimizing Dallas-Fort Worth Bus Transportation System Using Any Logic. *Journal of Sensors, IoT & Health Sciences*, 2(4), 40-55.
5. Oluwasammi, A., Aftab, M. U., Qin, Z., Ngo, S. T., Doan, T. V., Nguyen, S. B., ... & Nguyen, G. H. (2021). Features to text: a comprehensive survey of deep learning on semantic segmentation and image captioning. *Complexity*, 2021(1), 5538927.
6. Tan, F., Zhai, M., & Zhai, C. (2024). Foreign object detection in urban rail transit based on deep differentiation segmentation neural network. *Heliyon*, 10(17).
7. Xue, Z., Song, G., Guo, Q., Liu, B., Zong, Z., Liu, Y., & Luo, P. (2023). Raphael: Text-to-image generation via large mixture of diffusion paths. *Advances in Neural Information Processing Systems*, 36, 41693-41706.

8. Srinivasa Sai Abhijit Challapalli. (2024). Sentiment Analysis of the Twitter Dataset for the Prediction of Sentiments. *Journal of Sensors, IoT & Health Sciences*, 2(4), 1-15.
9. Yang, Z., & Yang, S. (2023). Multimedia image evaluation based on blockchain, visual communication design and color balance optimization. *Heliyon*, 9(12).
10. Huang, R., Zhang, Y., Liu, R., Song, Z., Ren, Z., & Pu, M. (2021). Application of Liver CT Image Based on Sueno Fuzzy C-Means Graph Cut and Genetic Algorithm in Feature Extraction and Classification of Liver Cancer. *Journal of Medical Imaging and Health Informatics*, 11(9), 2481-2489.
11. A.B. Hajira Be. (2024). Feature Selection and Classification with the Annealing Optimization Deep Learning for the Multi-Modal Image Processing. *Journal of Computer Allied Intelligence*, 2(3), 55-66.
12. Adegboye, O. R., Feda, A. K., Ishaya, M. M., Agyekum, E. B., Kim, K. C., Mbasso, W. F., & Kamel, S. (2023). Antenna S-parameter optimization based on golden sine mechanism based honey badger algorithm with tent chaos. *Heliyon*, 9(11).
13. Zhang, M., Wang, J., Cao, X., Xu, X., Zhou, J., & Chen, H. (2024). An integrated global and local thresholding method for segmenting blood vessels in angiography. *Heliyon*, 10(22).
14. Hollandi, R., Moshkov, N., Paavolainen, L., Tasnadi, E., Piccinini, F., & Horvath, P. (2022). Nucleus segmentation: towards automated solutions. *Trends in Cell Biology*, 32(4), 295-310.
15. Mencattini, A., Spalloni, A., Casti, P., Comes, M. C., Di Giuseppe, D., Antonelli, G., ... & Martinelli, E. (2021). NeuriTES. Monitoring neurite changes through transfer entropy and semantic segmentation in bright-field time-lapse microscopy. *Patterns*, 2(6).
16. Srinivasa Sai Abhijit Challapalli, Bala kandukuri, Hari Bandireddi, & Jahnvi Pudi. (2024). Profile Face Recognition and Classification Using Multi-Task Cascaded Convolutional Networks. *Journal of Computer Allied Intelligence*, 2(6), 65-78.
17. Bazhenov, E., Jarsky, I., Efimova, V., & Muravyov, S. (2024, September). EvoVec: Evolutionary Image Vectorization with Adaptive Curve Number and Color Gradients. In *International Conference on Parallel Problem Solving from Nature* (pp. 383-397). Cham: Springer Nature Switzerland.
18. Reddy, A. M., Reddy, K. S., Jayaram, M., Lakshmi, N. V. M., Aluvalu, R., Mahesh, T. R., ... & Alex, D. S. (2022). Research Article An Efficient Multilevel Thresholding Scheme for Heart Image Segmentation Using a Hybrid Generalized Adversarial Network.
19. Dong, Y., Fei, C., Zhao, G., Wang, L., Liu, Y., Liu, J., ... & Zhao, X. (2023). Registration method for infrared and visible image of sea surface vessels based on contour feature. *Heliyon*, 9(3).
20. Wang, G., Hu, J., Zhang, Y., Xiao, Z., Huang, M., He, Z., ... & Bai, Z. (2024). A modified U-Net convolutional neural network for segmenting periprosthetic adipose tissue based on contour feature learning. *Heliyon*, 10(3).
21. Javidan, S. M., Banakar, A., Vakilian, K. A., Ampatzidis, Y., & Rahnama, K. (2024). Early detection and spectral signature identification of tomato fungal diseases (*Alternaria alternata*, *Alternaria solani*, *Botrytis cinerea*, and *Fusarium oxysporum*) by RGB and hyperspectral image analysis and machine learning. *Heliyon*, 10(19).

22. Zhao, S., Yao, X., Yang, J., Jia, G., Ding, G., Chua, T. S., ... & Keutzer, K. (2021). Affective image content analysis: Two decades review and new perspectives. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(10), 6729-6751.
23. Panda, J., & Meher, S. (2024). Recent advances in 2d image upscaling: a comprehensive review. *SN Computer Science*, 5(6), 735.
24. Liu, Y., Sun, Y., Xue, B., Zhang, M., Yen, G. G., & Tan, K. C. (2021). A survey on evolutionary neural architecture search. *IEEE transactions on neural networks and learning systems*, 34(2), 550-570.
25. Plutino, A., Barricelli, B. R., Casiraghi, E., & Rizzi, A. (2021). Scoping review on automatic color equalization algorithm. *Journal of Electronic Imaging*, 30(2), 020901-020901.