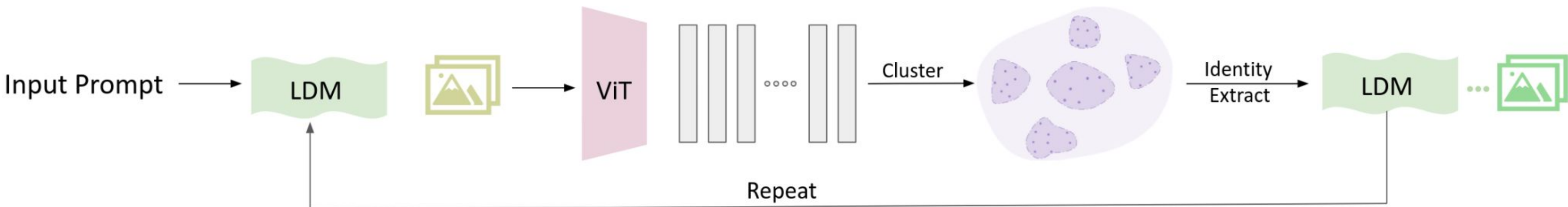
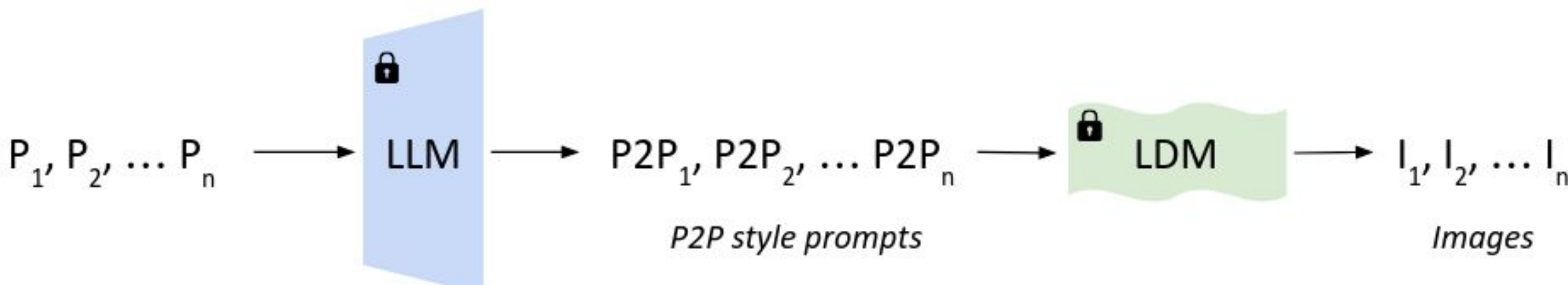
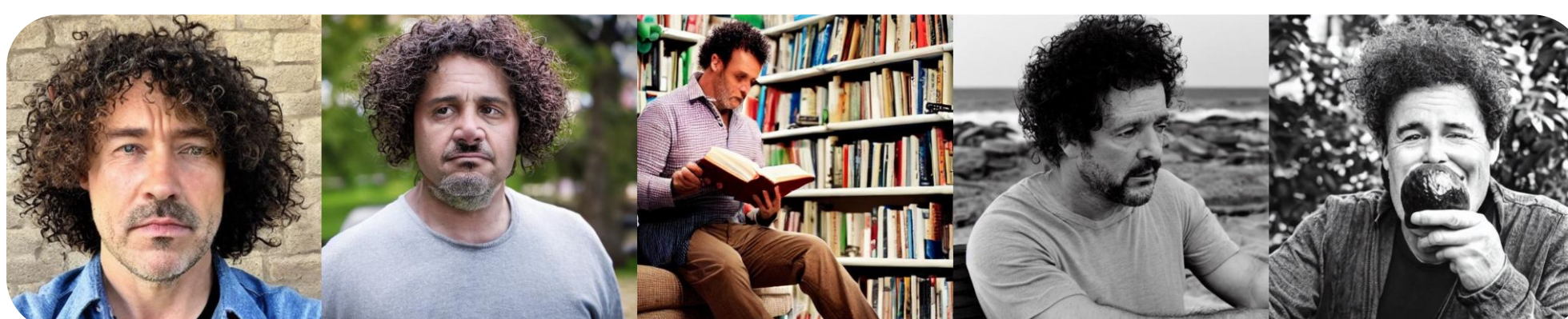
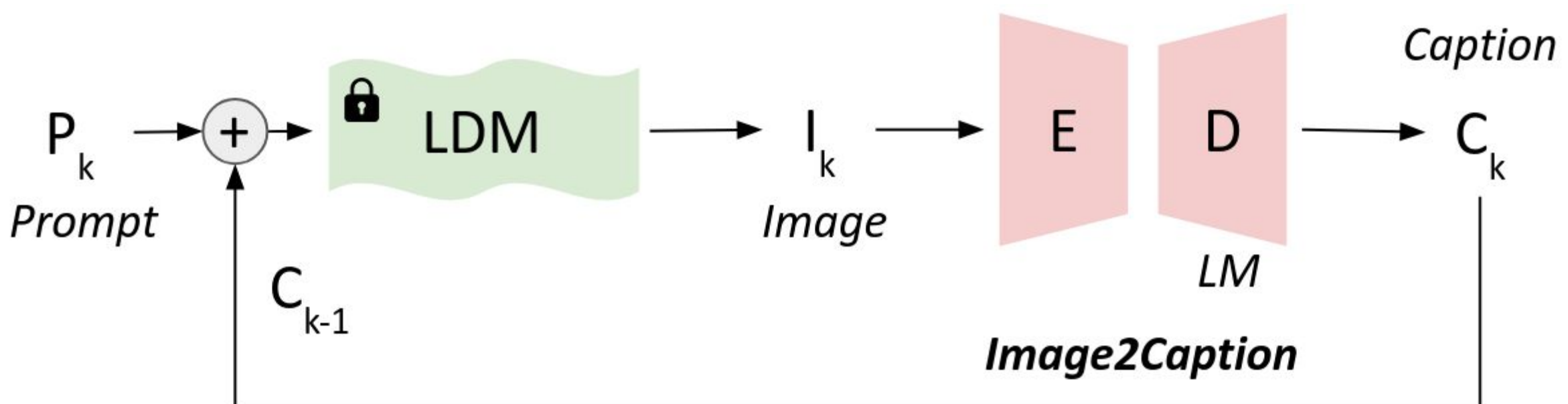
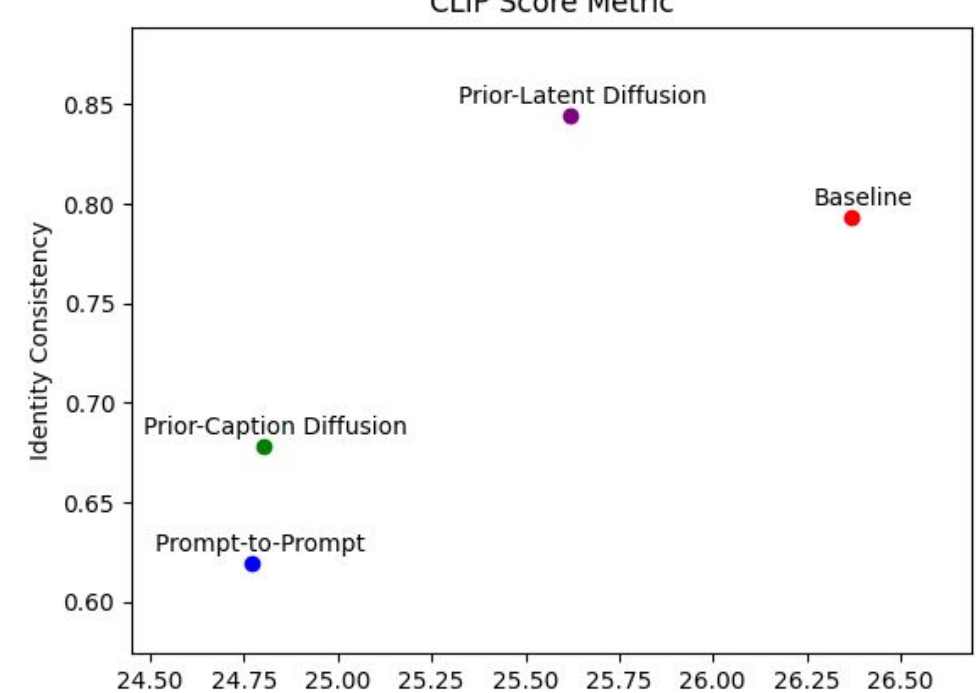
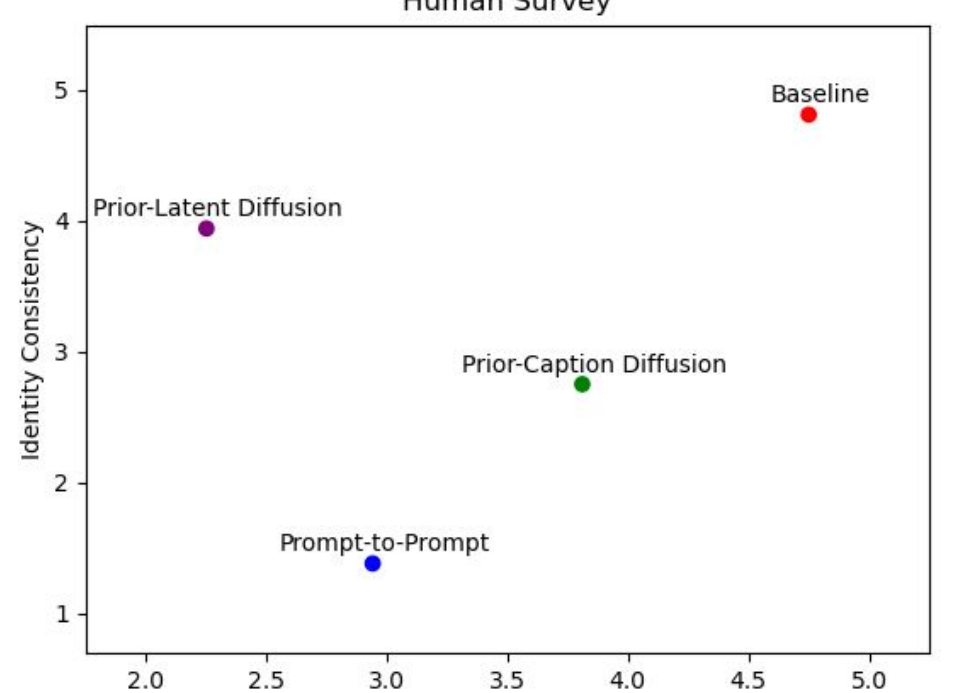
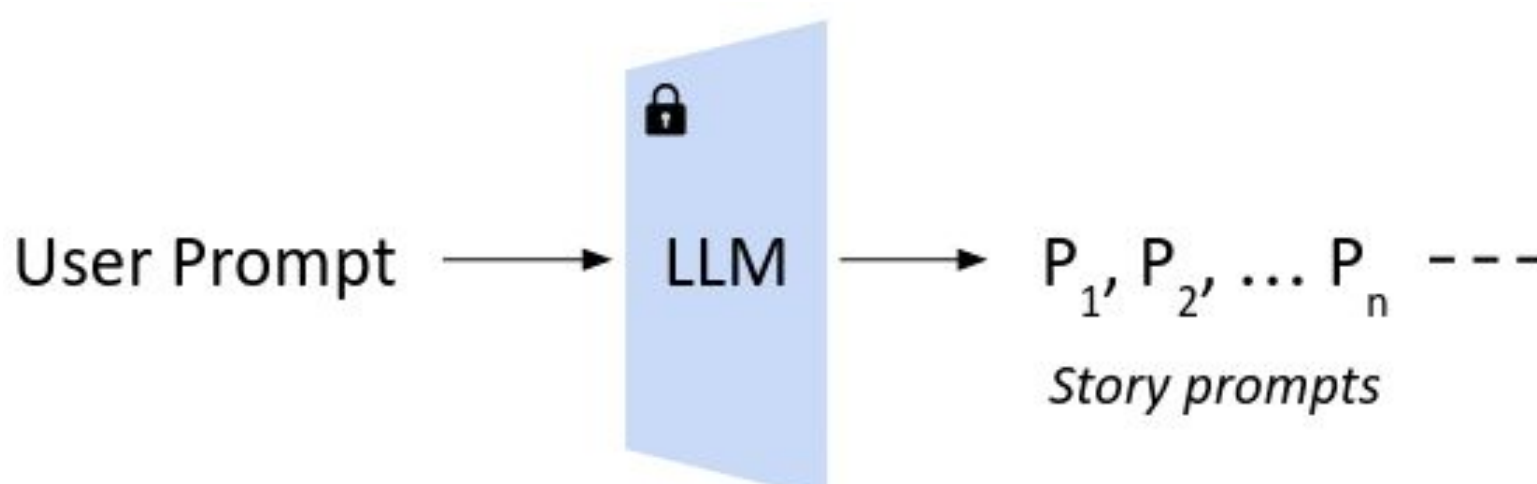
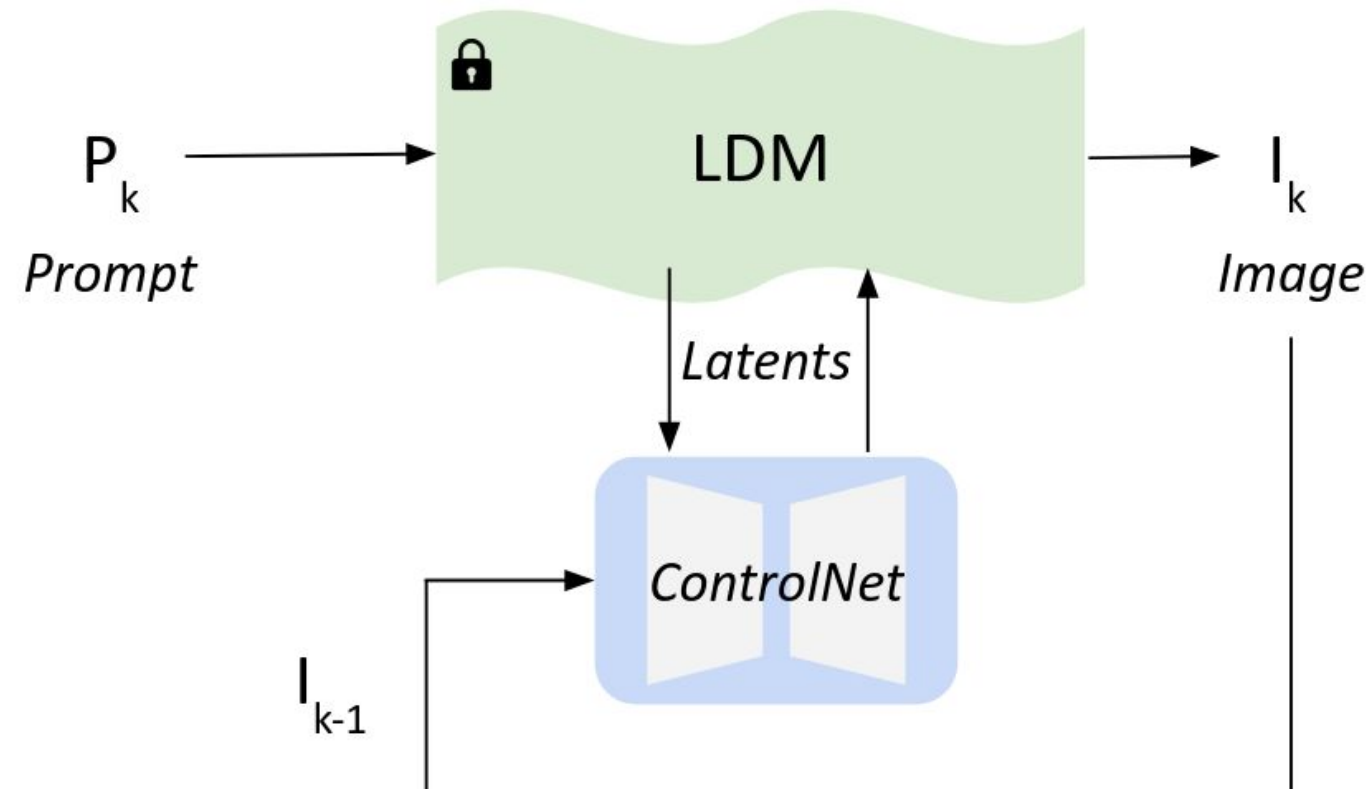


Introduction	Baseline: The Chosen One ²	Results						
This project explores novel methods for generating sequential images that form coherent visual storylines while maintaining consistent character representations. We evaluate three approaches: Prompt-to-Prompt: Progressive image generation using prior frames. Prior-Caption Diffusion: Iterative caption generation guiding subsequent images. Prior-Latent Diffusion: Latent space alignment for temporally consistent image generation. Our goal: Develop approaches that enhance storyline fidelity while ensuring visual consistency in character representation.	 Iterative refinement of generated images based on semantic clustering, guided by cohesive feature extraction and shared characteristics. Clustering is performed on image embeddings extracted using pre trained DINO v2, guiding the training of the LDM model towards consistent representations.	“A photo of a 50 years old man with curly hair...” in the park reading a book at the beach holding an avocado  Baseline						
Evaluation	Method 1: Dynamic Base Prompt-to-Prompt	Method 1						
Metrics: Character Consistency Text-Image Alignment <table><thead><tr><th>Human Survey</th><th>CLIP Model</th></tr></thead><tbody><tr><td>“Rate the main character's consistency across these 5 images”</td><td>Directional CLIP similarity (openai/clip-vit-large-patch14)</td></tr><tr><td>“Rate how accurately the image represents the prompt”</td><td>CLIP Score (openai/clip-vit-base-patch16)</td></tr></tbody></table> Evaluation Methods: CLIP Models: Compute quantitative values (Text-2-Image and Image-2-Image CLIP scores¹). Human Surveys: Gather subjective ratings for perceptual quality.	Human Survey	CLIP Model	“Rate the main character's consistency across these 5 images”	Directional CLIP similarity (openai/clip-vit-large-patch14)	“Rate how accurately the image represents the prompt”	CLIP Score (openai/clip-vit-base-patch16)	Sequential image generation conditioned on prior frames, with storyline prompts formatted for consistency. 	 Method 2
Human Survey	CLIP Model							
“Rate the main character's consistency across these 5 images”	Directional CLIP similarity (openai/clip-vit-large-patch14)							
“Rate how accurately the image represents the prompt”	CLIP Score (openai/clip-vit-base-patch16)							
	Method 2: Prior-Caption Diffusion	Method 2						
	Iterative image generation guided by captions extracted from prior images using a vision captioning³ + language model pipeline. 	 Method 3 CLIP Score Metric  Human Survey 						
User Input to Story Prompts	Method 3: Prior-Latent Diffusion	References						
	Leverages ControlNet⁴ to condition subsequent generations directly on previously generated images and the current story prompt. 	- Radford, Alec, et al. "Learning transferable visual models from natural language supervision." - Avrahami, Y.; Zada, D.; Fried, O.; Aberman, K. The Chosen One: Consistent Characters in Text-to-Image Diffusion Models. IEEE Trans. Pattern Anal. Mach. Intell. 2024, 46, 1234–1246. - Li, Junnan, et al. "Blip: Bootstrapping language-image pre-training for unified vision-language understanding and generation." International conference on machine learning. PMLR, 2022. - Zhang, Lvmin, Anyi Rao, and Maneesh Agrawala. "Adding conditional control to text-to-image diffusion models."						