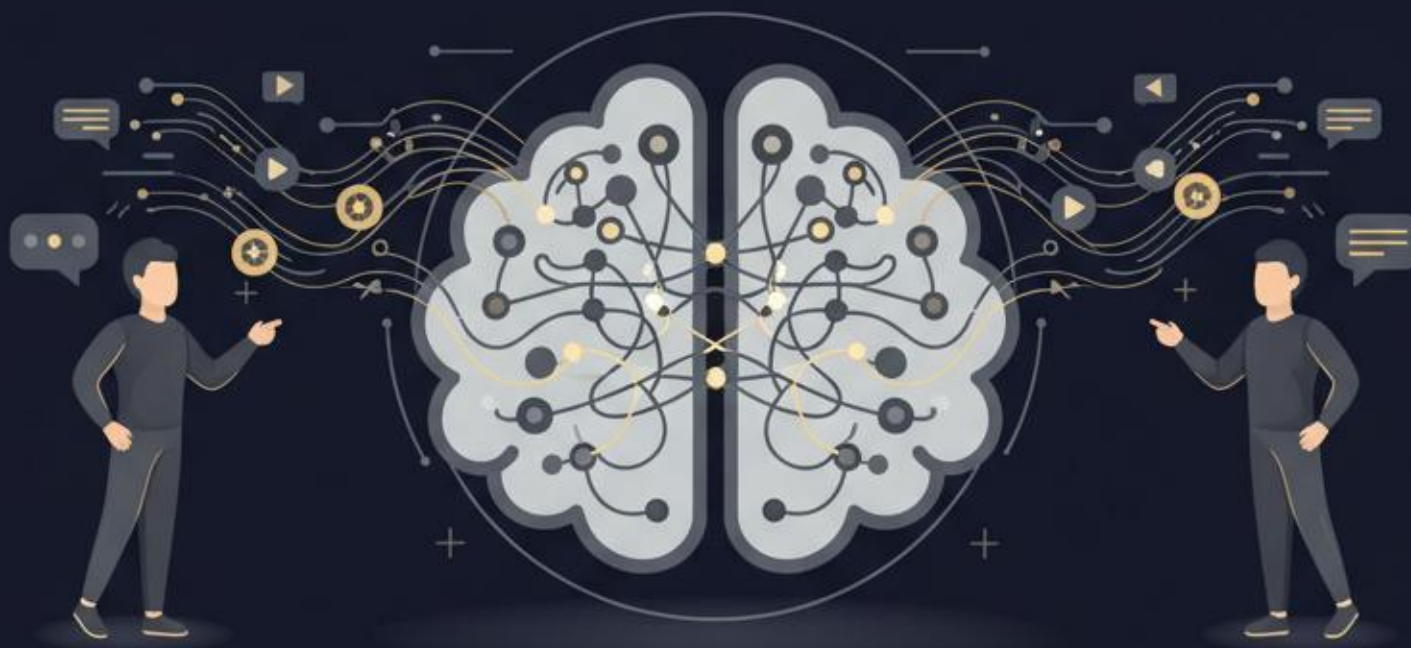


LEARNING CROSS-MODAL CONTEXT FOR VIDEO CAPTIONING

Cross-modal context interaction



Video captioning, which aims to automatically generate natural language descriptions for video content, is a challenging task due to the complex video data and the intricate relationship between video data and the natural language.



This document introduces a novel Multi-level Cross-modal Context (MCC) model designed to address the limitations of the existing video captioning approaches.



The MCC model characterized by its effectiveness in capturing and integrating multi-level cross-modal contexts between visual and textual modalities, leading to more accurate and diversified video descriptions.

AUTHORS / AFFILIATIONS

Problem Statement and Limitations of Existing Methods

Separate/ Superficial Interaction



- Visual & Textual features separately.
- Single, shallow interaction level

Limited Context Modeling



- Lack of fine-grained object detail.
- Struggle with global scene understanding.

Generic & Repetitive Captions

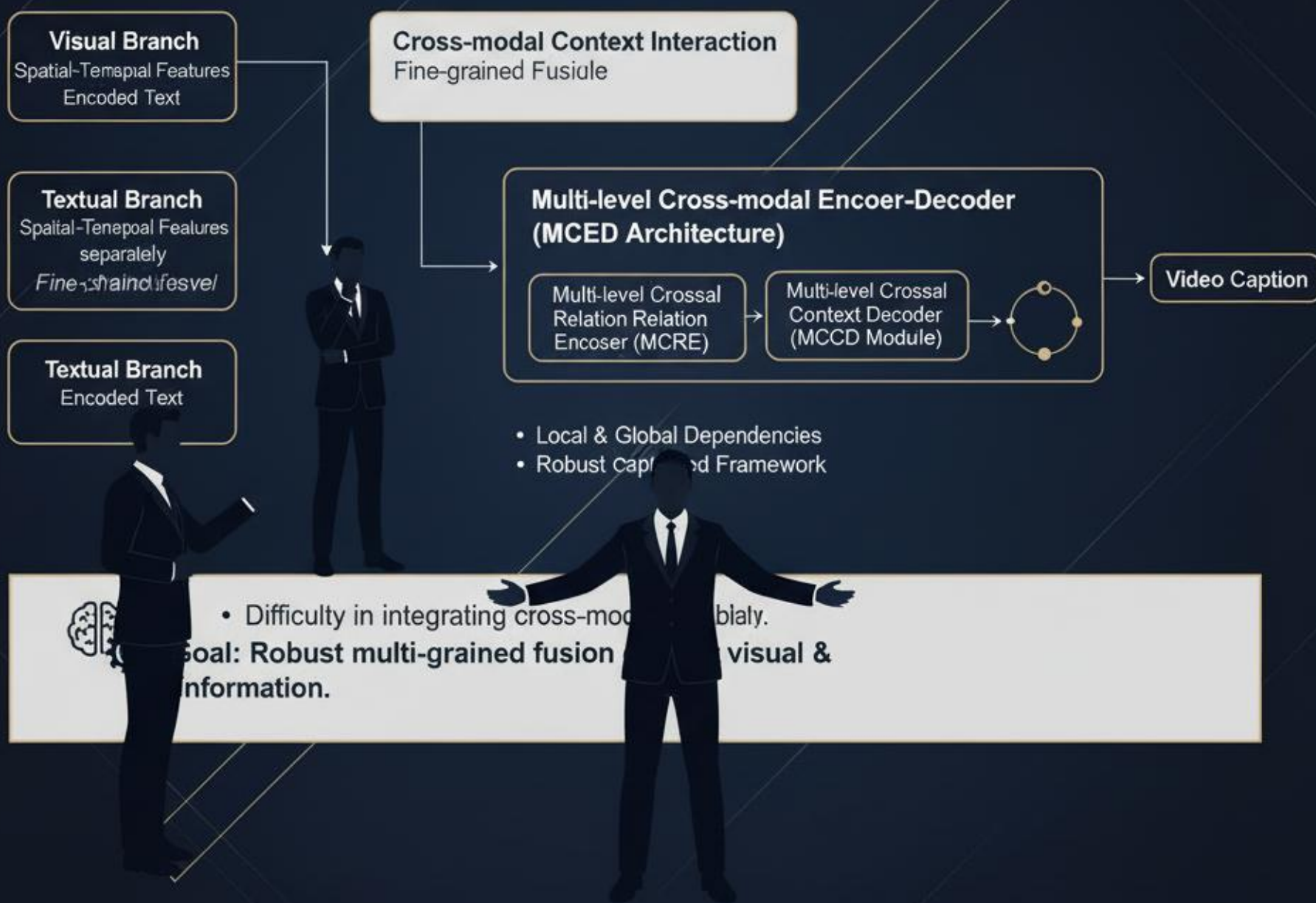


- Repetitive phrases.
- Lacks diversity & naturalness.



- Difficulty in integrating cross-modal context globally & locally.
- Goal: Robust multi-grained fusion of visual & textual information.**

Multi-level Cross-modal Context (MCC) Model Overview



Cross-modal Context Interaction (CCI) Module: Fine-grained Fusion



Module Overview

- Crucial for fine-grained fusion of visual & textual features. Establishes a relationship between visual & textual interaction



Process & Output

- Visual features V and textual embedding E(T) input. Uses Attention Mechanism



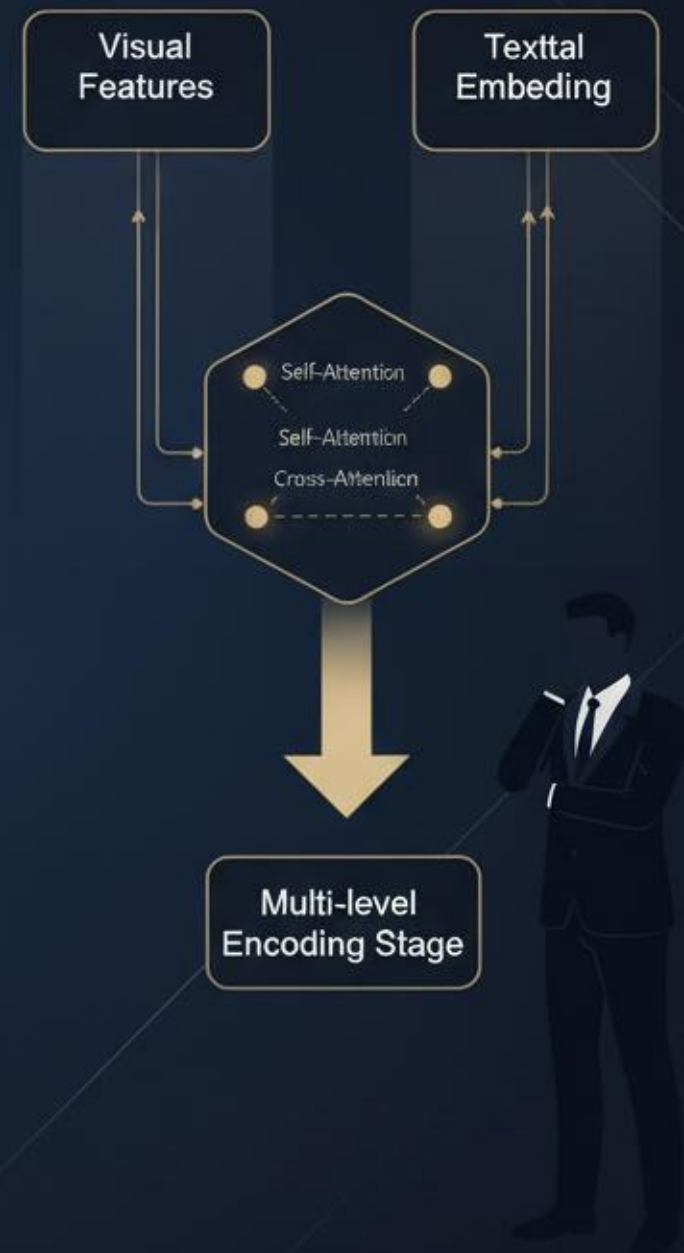
Series self-attention & cross-attention layers
Each modality informs & refines the other
Captures dynamic relationships at granular level

- Computes attention weights for cross-modal CCI at granular level.



Mathematical Foundation

$$Q = \frac{\text{Attention}(Q, K, V)}{QK^T} = \frac{\text{softmax}_k \left[\frac{QK^T}{\sqrt{d_k}} \right] V}{\sqrt{d_k}}$$



Multi-level Cross-modal Encoder-Decoder (MCED) Architecture

The Multi-level Cross-modal Relation architecture is responsible for processing the fused cross-modal contexts and generating the video captions. This has three main main components: Multi-level Cross-modal Encoder (MCRE), and other an encoder, and textual components:

1. Multi-level Cross-modal Relation Encoder (MCRE)



- Leverages interactive context from Builds rich, multi-grained video representations
- Uses multiple encoder layers for semantic levels.
- Self-attention for intra-attention for for inter-modal
- Hierarchical encoding for global-details.

2. Multi-level Cross-modal Context Decoder (MCCD)



- Generates caption word-vectors-Conditioned on limited on MCRE representations.
- Dynamic attention selects relevant co-variant info from different levels and layers with
- masked self-rooted self-attention & cross-attention.

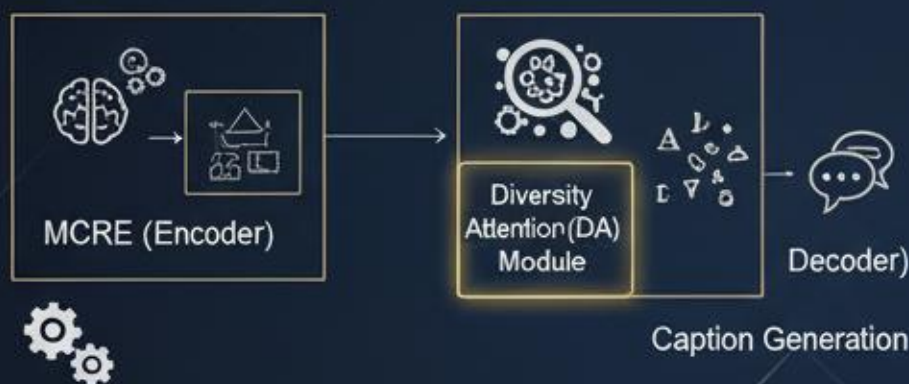
Multi-level Cross-modal Encorer-Decoder (MCED)

Diversity Attention (Cont):

3. Diversity Attention: Fine-grained Fusion

3. Diversity Attention (DA) Module

MCED Architecture



- Dynamically adjusts attention weights.
- Encourage diversity in word choice.
- Prevents repetitive phrases.

- Modulates standard attention with diversity component (DA) word choice.
(λ : diversity emphasis hyperparameter).



- Difficulty in integrating attentely

Goal: multi-grained visual fusion: $s'_{\text{attn}} + s_{\text{attn}} = \lambda s^{\text{attn}}_{\text{div}} \}$
(λ : diversity emphasis hyperparameter).



Experimental Setup: Datasets and Evaluation Metrics

MSVD Dataset



- 1,970 YouTube clips
42 hours total
- 40 sentences/video (avg.).
Split: 1200 Train, 100 Val, 670 Test
- Diverse content, high, high
annotation variability

MSR-VTT Dataset



- 1,0000 video clips
- 20 English sentences/video
Split: 6513 Train, 497 Val Test
- Wide categories: people, sports,
Complex, longer
descriptions

Evaluation Metrics



- CIDEr (Semantic similarity)
- METEOR (N-gram, fluency,
(Longest-L (Longest Common
Subsequence
- BLEU@1-4 (Ngram overlap)



Measures quality: n-gram overlap, semantic similarity, fluency

Main Results and Comparison on MSVD and MSR-VTT

MSVD Dataset

Metric	MSVD	METEOR	Previous SOTA	Bprevious
CIDer	98.2%	95.2%	31.1%	43.2%
METEOR	35.3%	35.7%	34.1%	45.3%
ROUGE-L	59.3%	59.3%	57.78	48.2%
BLEU@4	45.1%	56.2%	56.2%	43.2%

MSR-VTT Basaeet

Model	CIDer	METEOR	S2.2-L	MSU-VT
MCC Model	47.6%	31.1%	52.2%	31.6%
Transformer	44.2	38.4	28.5	28.0
HCRN	45.8	28.8	50.9	28.1
X-LAN	46.5	20.9	51.11	30.5

- The MCC model consistently achieved the of-state or highly compettlince across both MSVD globally, demonisrting accurate and diverse video capbally.



The MCC model consistently achieved strørte of uretative baselines... our MCC model our MSV-VTT deriT datesets, to model suprases X-LAN by **+1.2%** in **METEOR** and demonsating its effecteuemess in significantly Transferrer by **+2.7%** in **METOR** and **+3.4%** in **3.4%** in CIDE.

Ablation Study: Impact of Key Components

Component Removed

	CIDer (%) Removed	METEOR (%)	Key Insight
Baseline (w/o MCED)	43.1	Critical role of multi-processing	Repetitive phrases.
w/o MCRE	27.8	Encoder builds rich cross-modal	Encoder levention
w/o MCCD	45.7	cross-modal representations	Lacks caption diversity & quality
w/o MDA	47.0	Decoder leverages multi-level context	Lacks diversity & quality
Full MCC Model:	47.6	State-of-state performance	



Each component is vital for ineris & locally.
Each component for MCC model's superior performance

Qualitative Analysis and Diversity Improvement

MCC Model: Richer Descriptions



More Accurate & Diverse: Captures specific details.



Baseline: "A person is playing guitar."



MCC: "A man strumming acoustic guitar, focusing intently on the music."

Diversity Attention (DA) Module



Encourages Broader
Avoids repetition



Encourages Broader Vocabulary:



Dynamic Attention Adjustment:
Balances accuracy & variation.



More Natural & Human-like

Figure 3: Qualitative Improvements



BASELINE: "A person is playing guitar."

DIVERSITY
→



MCC + DA: "A man is acoustic guitar, focusing on the music."

Conclusion and Future Work

Key Contributions

- 1  Cross-modal Context Interaction (CCI): Enables fine-grained fusion of visual and textual features. 
- 2  Multi-level Cross-modal Encoder-Decoder (MCED): it Comparing MCRE and MCCD, it processes multiple semantic levels for robust state CIDEr on 97.6% for robust encoding and decoding.

Future Work

- 
-  Extend to incorporate external knowledge for richer semantic understanding.
 -  Explore sophisticated spatiotemporal reasoning (e.g. graph-based representations) video content).



Experimental results on MSVD and MSR TT datasets demonstrated demonstrate to performance, achieving state up MS.2% CIDEr on MSER and MSR CIDEr. Ablation studies confirmed contributions confirmed MCC on each critical component by produce accurate, detailed, and diverse descriptions.