

附件 1

南开大学推荐免试攻读研究生申请表暨诚信承诺书

学号	2312220	姓名	宋昊谦	性别	男
民族	汉	出生日期	2004 年 11 月 27 日	政治面貌	共青团员
学院	密码与网络空间安全学院		专业	密码科学与技术	
手机号码	16629057198		电子邮箱	2312220@mail.nankai.edu.cn	
申请类型 (单选)	☒ 普通类推免名额 ☐ 支教团推免名额				
本人与学院遴选指标相关的材料列表	(学生可将本人与本学院遴选指标相关的材料列在下方，如科研成果、竞赛获奖、科研训练表现、参军入伍服兵役、志愿服务、国际组织实习等) 1. ICASSP 2026 第一作者论文，CCF-B 类 2. 全国大学生数学建模竞赛天津赛区，二等奖 3. 社会实践 2 次 4. 志愿服务时长共 25.5 小时 5. 学院迎新晚会 1 次				
本人充分了解并理解《南开大学推荐优秀应届本科毕业生免试攻读研究生工作实施办法》(南党发〔2021〕63 号)和所在学院《推免细则》，并郑重承诺： 1. 本人所提供的一切材料真实可靠，无弄虚作假，如有不实之处，本人愿意承担由此造成的一切后果； 2. 如获得推免生资格，本人将慎重填报志愿，绝不浪费学校推免名额。 承诺人： 年 月 日					

(本申请书必须由申请人亲笔签名方为有效，交所在学院本科教学工作办公室备查)

综合素质评价审查材料目录（不包含公能素质部分）

学院：密码与网络空间安全学院 专业：密码科学与技术 学号：2312220 姓名：宋昊谦

请根据个人实际情况填写此目录（其他项目可自行添加行），并按填写顺序排序所有证明材料
后附证明材料要求真实、直观。
此目录纸张横向单面打印、附在“承诺书”后

		加分项目 ① 赛事获奖请直接直接填写右边各列 ② 论文加分请写明论文题目后填写右边各列 ③ 辅修专业课程请写明辅修专业 ④ 参军入伍及系列情况写“参军入伍” 每类事项行数不够请在该类目下自行添加	此处填写获奖等级： ① 赛事获奖请填写“一等奖/二等奖/三等奖/金奖/银奖/铜奖” ② 论文发表请填写“发表期刊-类别”（eg. “IEEE TDSC”-CCF 英文 A 类） ③ 辅修专业课程总学分数 ④ 参军入伍/嘉奖/优秀士兵/三等功及以上	证明材料类别： ① 赛事获奖证书/获奖公告页面截图 （材料里需标注出个人获奖信息所在行） ② 录用邮件页面截图/期刊网页截图+论文原文文本（证明自己是一作） ③ 辅修成绩单
特殊学术专长	重要赛事	高教社杯全国大学生数学建模竞赛本科组		
		全国大学生电子设计竞赛本科组		
		“挑战杯”全国大学生课外学术科技作品竞赛		
		中国国际大学生创新大赛（原中国国际“互联网+”大学生创新创业大赛）		
科研训练	各专业代表性赛事	ACM-ICPC 亚洲区域赛		
		全国大学生信息安全竞赛		
		全国密码技术竞赛		
		全国大学生计算机系统能力大赛		

		全国大学生物联网设计竞赛		
	一般赛事	全国大学生数学建模竞赛天津赛区	二等奖	赛事获奖证书
		天津市“挑战杯”大学生课外学术科技作品竞赛		
		中国国际大学生创新大赛天津赛区(原中国国际“互联网+”大学生创新创业大赛天津赛区)		
		天津市物联网大赛		
		京津冀大学生信息安全网络攻防大赛(原天津市大学生信息安全网络攻防大赛)		
	辅修			
特殊学术专长	论文发表	DUAL-PATH COMPRESSION FOR REAL-TIME MULTIMODAL CLICKBAIT DETECTION:QUANTIZATION AND DISTILLATION	ICASSP 2026-CCF 英文 B 类	会议论文集网页截图，录用邮件截图、论文全文
	论文发表			
	论文发表			
参军入伍				
其他材料				



证书编号: TJ-2025-SXJM-0182B

全国大学生数学建模竞赛 获奖证书

南开大学

宋昊谦、赵甫文、姜若然 同学:

在2025年全国大学生数学建模竞赛天津赛区竞赛中,
您的作品获得本科组二等奖。

指导教师: 李忠华
特发此证, 以资鼓励。

天津市教育委员会

2025年11月18日

中国工业与应用数学学会

IEEE Xplore®

Browse ▾

My Settings ▾

Help ▾

Access provided by:
NANKAI UNIVERSITY

Sign Out

IEEE

All

Q

ADVANCED SEARCH

Conferences > ICASSP 2026 - 2026 IEEE Inter... ?

Dual-Path Compression for real-time Multimodal Clickbait Detection: Quantization and Distillation

Publisher: IEEE

Cite This

PDF

Haoqian Song

Haoran Yin

Fuwen Zhao

Yiran Du

Xuan Luo

Xindi Ma

All Authors

82

Full

Text Views

R

©

More Like This

Implementation of Grey Scale Normalization in Machine Learning & Artificial Intelligence for Bioinformatics using Convolutional Neural Networks

2021 6th International Conference on Inventive Computation Technologies (ICICT)

Published: 2021

Machine learning and deep neural network — Artificial intelligence core for lab and real-world test and validation for ADAS and autonomous vehicles: AI for efficient and quality te...

Abstract

Document Sections

1. INTRODUCTION

2. PROPOSED METHOD

Abstract:

While multimodal pre-trained models achieve state-of-the-art performance in clickbait detection, their computational requirements hinder practical deployment. We present a dual-path compression framework for multimodal clickbait detection, evaluated on a new fine-grained Chinese dataset (15,012 samples) that distinguishes three deception mechanisms: non-clickbait, curiosity gap, and content mismatch. Our approach employs INT8 quantization-aware training for server deployment and knowledge distillation for edge devices, achieving complementary compression-

ICASSP 2026: Review Results [Paper #9628]

haoqiansong@qq.com

ICASSP 2026

papers@2026.ieeeicassp.org

发至 3213152552、2313547、Fuwen Zhao、Yiran Du、Xuan Luo、Xindi Ma、Chaoyuan Zuo 收起

2026/01/18 01:41

发件人: ICASSP 2026<papers@2026.ieeeicassp.org>

收件人: 3213152552<3213152552@qq.com>; 2313547<2313547@mail.nankai.edu.cn>; Fuwen Zhao<zfw hue@163.com>; Yiran Du<2311390@mail.nankai.edu.cn>; Xuan Luo<2312615@mail.nankai.edu.cn>; Xindi Ma<2213369@mail.nankai.edu.cn>; Chaoyuan Zuo<zuocy@nankai.edu.cn>;

主题: ICASSP 2026: Review Results [Paper #9628]

时间: 2026年01月18日 (周日) 01:41

↩

↶

↷

目 备注

...

Dear Haoqian Song, Haoran Yin, Fuwen Zhao, Yiran Du, Xuan Luo, Xindi Ma, Chaoyuan Zuo,

Paper ID: 9628

Title: DUAL-PATH COMPRESSION FOR REAL-TIME MULTIMODAL CLICKBAIT DETECTION: QUANTIZATION AND DISTILLATION

Congratulations!

We are pleased to inform you that the above-listed manuscript has been accepted for presentation at IEEE ICASSP 2026 (<https://2026.ieeeicassp.org>), which will be held as an in-person event in Barcelona, Spain, from May 4 to 8, 2026.

Over the past few months, the IEEE Signal Processing Society's Technical Committees worked very hard reviewing the 11,120 papers submitted to ICASSP 2026. This was the largest and most comprehensive ICASSP in IEEE SPS history. Expert reviewers evaluated submissions in their areas of expertise, resulting in a highly competitive selection process with an acceptance rate of about 40%.

Thank you for contributing to IEEE ICASSP 2026! To celebrate your paper acceptance, we've created a shareable graphic you can use to promote your work and invite your network to attend. Use this link to spread the word across LinkedIn, X, Slack, WhatsApp, email, and more: <https://app.gleanin.com/share/c/43733>

DUAL-PATH COMPRESSION FOR REAL-TIME MULTIMODAL CLICKBAIT DETECTION: QUANTIZATION AND DISTILLATION

Haoqian Song^{2,†}, Haoran Yin^{2,†}, Fuwen Zhao^{2,†}, Yiran Du¹, Xuan Luo¹, Xindi Ma³, Chaoyuan Zuo^{1,*}

¹School of Information and Communication, ²College of Cryptology and Cyber Science,

³School of Literature, Nankai University, Tianjin, China

ABSTRACT

While multimodal pre-trained models achieve state-of-the-art performance in clickbait detection, their computational requirements hinder practical deployment. We present a dual-path compression framework for multimodal clickbait detection, evaluated on a new fine-grained Chinese dataset (15,012 samples) that distinguishes three deception mechanisms: non-clickbait, curiosity gap, and content mismatch. Our approach employs INT8 quantization-aware training for server deployment and knowledge distillation for edge devices, achieving complementary compression-efficiency trade-offs. We conduct comprehensive ablation studies examining modality contributions, compression strategies, and module-level sensitivity to guide deployment decisions. Results reveal that while both methods maintain strong overall accuracy (93.3% for INT8, 88.2% for distilled), they struggle with minority class detection, particularly curiosity gap categorization. These findings provide practical deployment strategies for resource-constrained scenarios and highlight the need for compression-aware architectures in multimodal systems.

Index Terms— Multimodal compression, quantization-aware training, knowledge distillation, real-time inference, clickbait detection

1. INTRODUCTION

Clickbait content has emerged as a widespread challenge in digital media ecosystems worldwide, with distinct patterns emerging across different linguistic and cultural contexts [1, 2]. In the Chinese digital ecosystem, platforms like Toutiao, WeChat, and Weibo serve over a billion users daily, where clickbait employs distinctive techniques ranging from four-character idiom manipulations to culturally specific visual metaphors [3, 4]. Recent advancements in detection methods have shifted from feature-based approaches to large-scale multimodal pre-trained language models (PLMs) that jointly analyze text and images [5]. While these models, often combining powerful text encoders with visual models and

[†]Equal contribution.

^{*}Corresponding author.

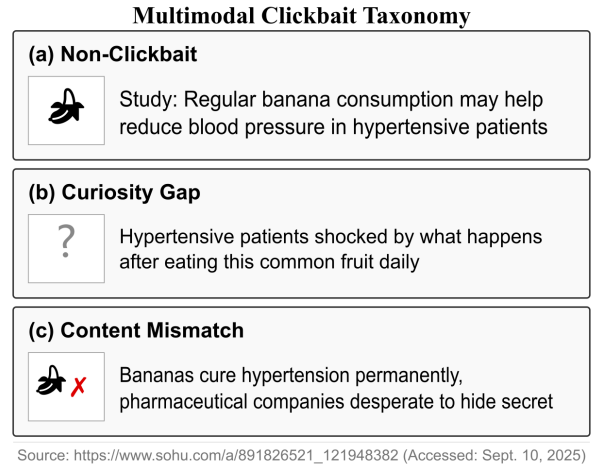


Fig. 1: Three-class clickbait taxonomy with examples.

syntactic analyzers like Graph Attention Networks (GAT) [6], achieve state-of-the-art accuracy, their hundreds of millions of parameters create a critical “deployment gap” for real-time content moderation systems.

Model compression techniques offer a promising solution to this computational challenge, with recent advances in quantization [7, 8] and distillation [9, 10] demonstrating success on English models. However, applying these techniques to Chinese multimodal architectures presents unique challenges: the compression must preserve both character-level semantics and cross-modal alignment while handling culturally specific patterns. Moreover, current detection systems typically employ binary classification, failing to distinguish between fundamentally different deception mechanisms. Drawing from information gap theory [11] and curiosity-driven behavior [12], we propose a fine-grained three-class taxonomy illustrated in Figure 1: (1) non-clickbait with factual headlines, (2) curiosity gap that exploits psychological curiosity through strategic information withholding without factual misrepresentation [13], and (3) content mismatch involving deliberate false claims.

This paper presents a comprehensive framework that addresses both the deployment efficiency and classification

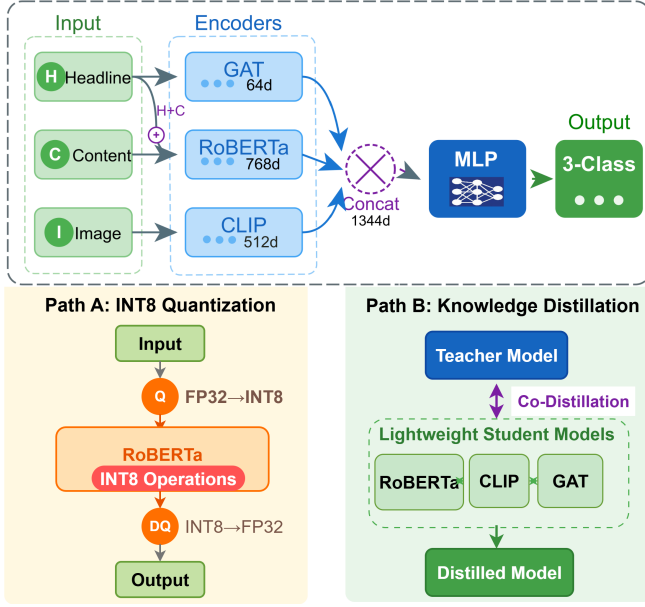


Fig. 2: Overview of our compression framework. Top: baseline multimodal architecture for three-class clickbait detection. Bottom: two compression paths - (A) INT8 quantization applied to the text encoder only, maintaining original architecture; (B) knowledge distillation creating lightweight student models.

granularity challenges. Building on recent multimodal architectures [5], our main contributions are: (1) We develop a fine-grained three-class taxonomy for Chinese clickbait detection, validated on a newly annotated dataset of 15,012 samples. (2) We evaluate two compression strategies on Chinese multimodal models, with INT8 quantization-aware training achieving $8.1\times$ compression at 93.3% accuracy for server deployment and knowledge distillation providing $12.6\times$ compression at 88.2% accuracy for edge devices. (3) We uncover differential compression sensitivity across clickbait types, revealing that different deception mechanisms rely on distinct model capabilities, which enables targeted deployment strategies for real-world applications.

2. PROPOSED METHOD

Figure 2 presents our dual-path compression framework, which deploys multimodal clickbait detectors through INT8 quantization and knowledge distillation.

2.1. Baseline Multimodal Architecture

Our baseline employs a prompt-tuning paradigm integrating three complementary encoders to capture different aspects of clickbait patterns.

Semantic Encoding. We use Chinese-RoBERTa-wwm-ext [14] with learnable continuous prompts $\mathcal{T}(\cdot)$ wrapping the

input headlines. The model produces a 768-dimensional representation $\mathbf{h}_{\text{text}}$ from the [CLS] token, with a verbalizer $\mathcal{V}$ mapping predictions to class labels.

Syntactic Encoding. To capture grammatical ambiguities common in Chinese clickbait, we construct dependency graphs $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ using LTP [15] and encode them via GAT:

$$\mathbf{h}_{\text{gat}} = \text{GAT}(\mathcal{G}) \in \mathbb{R}^{64} \quad (1)$$

Visual Encoding. Images are encoded using CLIP ViT-B/32 [16], producing a 512-dimensional feature vector $\mathbf{h}_{\text{img}} \in \mathbb{R}^{512}$.

Multimodal Fusion. The three representations are concatenated and classified via a two-layer Multi-Layer Perceptron (MLP):

$$\mathbf{h} = [\mathbf{h}_{\text{text}}; \mathbf{h}_{\text{img}}; \mathbf{h}_{\text{gat}}] \in \mathbb{R}^{1344} \quad (2)$$

$$\hat{\mathbf{y}} = \text{Softmax}(\mathbf{W}_2(\text{ReLU}(\mathbf{W}_1\mathbf{h} + \mathbf{b}_1)) + \mathbf{b}_2) \quad (3)$$

We employ class-weighted cross-entropy loss to handle the imbalanced distribution described in Section 3.

2.2. INT8 Quantization-Aware Training

To accelerate inference and reduce the model footprint, we apply quantization-aware training (QAT) to the baseline’s text encoder, which is the main computational bottleneck. Unlike post-training quantization that often degrades accuracy, QAT simulates quantization effects during training, allowing the model to adapt to the reduced precision while preserving performance. Our implementation details are as follows:

Scope and Backend: We insert quantization operators (QuantStub/DeQuantStub) around the encoder inputs, multi-head attention blocks, and transformer outputs. For numerical stability, LayerNorm and Softmax operations remain in FP32. We adopt the `fbgemm` [17] backend, which is optimized for x86 CPU inference.

Training and Conversion: Starting from the FP32 baseline model, we perform QAT to allow the model to adapt to low-precision arithmetic. The training procedure follows the same hyperparameter settings as our baseline model. After training, the model is converted to a native INT8 format for deployment.

2.3. Knowledge Distillation for Lightweight Models

For edge device deployment, we construct compact distilled models that mirror the baseline’s three-encoder architecture with reduced dimensions:

Semantic Encoding: 4 layers with hidden dimension 384 and 6 attention heads, producing $\mathbf{h}_{\text{text}}^{(s)} \in \mathbb{R}^{384}$ projected to $\mathbb{R}^{768}$.

Visual Encoding: 6 layers with hidden dimension 256 and 4 attention heads, yielding $\mathbf{h}_{\text{img}}^{(s)}$ projected to $\mathbb{R}^{512}$.

Syntactic Encoding: Single layer with hidden dimension 64 and 2 attention heads, generating $\mathbf{h}_{\text{gat}}^{(s)} \in \mathbb{R}^{64}$.

We optimize students through multi-objective distillation:

$$\mathcal{L} = \alpha \mathcal{L}_{\text{CE}} + \beta \mathcal{L}_{\text{KD}} + \gamma \mathcal{L}_{\text{feat}} \quad (4)$$

where $\mathcal{L}_{\text{CE}}$ is cross-entropy loss on ground-truth labels, $\mathcal{L}_{\text{KD}}$ transfers teacher knowledge via KL divergence between softened probability distributions:

$$\mathcal{L}_{\text{KD}} = \tau^2 \cdot \text{KL}(\text{softmax}(z_t/\tau) \parallel \text{softmax}(z_s/\tau)) \quad (5)$$

and $\mathcal{L}_{\text{feat}}$ aligns intermediate representations:

$$\mathcal{L}_{\text{feat}} = \lambda_{\text{text}} \|\phi_t - \phi_s\|_2^2 + \lambda_{\text{img}} (1 - \cos(\psi_t, \psi_s)) + \lambda_{\text{gat}} \|\gamma_t - \gamma_s\|_2^2 \quad (6)$$

where ϕ, ψ, γ denote projected text, image, and GAT features. We employ L2 distance for text and syntactic features while using cosine similarity for visual features to respect their embedding space properties.

3. EXPERIMENTAL SETUP

3.1. Dataset

Existing Chinese clickbait datasets are either proprietary or limited to binary classification, which obscures important distinctions between deception mechanisms. We construct a new dataset from Toutiao, China’s largest news aggregation platform, by crawling news articles and their associated meta-data. The dataset¹ contains 15,012 samples annotated using the three-class taxonomy illustrated in Figure 1.

Three annotators independently labeled each sample, with final labels determined by majority voting (Cohen’s Kappa = 0.83). The corpus exhibits typical characteristics of Chinese online news: headlines average 27.1 characters, articles average 1,205 characters, with 1.85 images per article. The distribution shows class imbalance: 9,919 non-clickbait (66.1%), 3,931 curiosity gap (26.2%), and 1,162 content mismatch (7.7%). To address this imbalance, we apply weighted cross-entropy loss with weights inversely proportional to class frequencies: $w_0 = 0.50$, $w_1 = 1.27$, $w_2 = 4.31$.

3.2. Implementation Details

All models are trained for 5 epochs using Adam optimizer with learning rate 2×10^{-5} , batch size 32, and maximum sequence length 192. Knowledge distillation employs two-stage training: component-wise pre-training followed by joint optimization with composite loss weighting ($\alpha = 0.6$, $\beta = 0.3$, $\gamma = 0.1$) and temperature $\tau = 2$ to ensure stable convergence.

We conduct experiments on complementary hardware platforms: NVIDIA RTX 4090 GPU for FP32 baseline and distilled models, and Intel i9-12900K CPU with fbgemm

backend for INT8 quantized models. This dual-platform setup captures the distinct computational constraints of server-side and edge deployments.

3.3. Evaluation Methodology

We employ 10-fold cross-validation with stratified 80/20 splits to ensure robust evaluation across the imbalanced class distribution. Efficiency metrics encompass throughput (samples/second), inference latency (ms/sample), and memory footprint quantified for both GPU VRAM and CPU RAM. For effectiveness assessment, we report accuracy, macro F1-score, and per-class precision/recall, with macro F1 serving as the primary metric to equally weight minority categories.

We conduct ablation studies to understand compression behaviors: (1) modality contributions by selectively removing visual or syntactic encoders, (2) quantization strategies comparing per-tensor versus per-channel granularity, (3) distillation variants contrasting synergistic multi-objective optimization against basic KL divergence, and (4) module-level sensitivity analysis through layer-wise quantization.

4. RESULTS AND ANALYSIS

Table 1 reveals fundamental trade-offs between compression approaches. INT8 quantization maintains architectural fidelity (270.4M parameters) while achieving 8.1× storage reduction, preserving 93.3% accuracy with modest CPU speedup (1.4× over FP32). Knowledge distillation reduces parameters by 92% (21.5M), enabling 4.8× faster CPU inference at 19.7 samples/s, but sacrifices 10.6 percentage points in accuracy. This contrast demonstrates that architectural simplification yields greater runtime benefits than numerical optimization for edge deployment.

Figure 3 exposes severe class imbalance in compression impacts. While non-clickbait detection remains robust across all models, minority classes suffer disproportionately. Curiosity gap detection degrades from 95.2% (baseline) to 77.9% (quantized) and 58.2% (distilled), suggesting these subtle manipulation patterns require representational capacity that compression eliminates. Content mismatch detection shows similar vulnerability, indicating semantic contradiction detection relies on complex cross-attention mechanisms that neither method preserves. The quantized model maintains higher recall (91.7%) than precision (85.2%), while the distilled model exhibits balanced degradation (P: 88.0%, R: 87.2%), making it suitable for binary screening but inadequate for fine-grained categorization.

Ablation Analysis. Compression strategy choices significantly impact performance. Per-tensor quantization outperforms per-channel by 1.8 points (92.3% vs. 90.5%), as consistent scaling better preserves distributed representations. More dramatically, synergistic distillation surpasses basic KL diver-

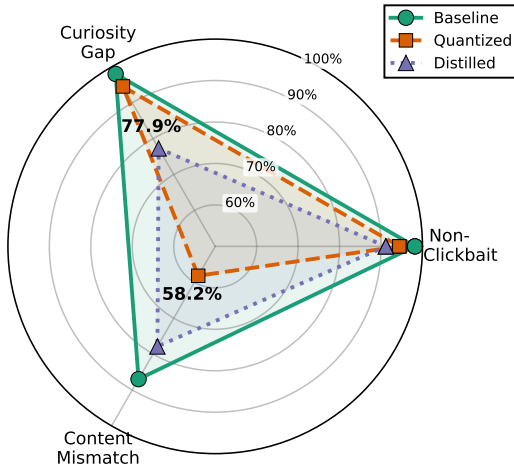
¹The dataset and code are available at <https://github.com/shq3526/DPCMCD>.

Table 1: Comprehensive evaluation of compression techniques and ablation studies.

(a) Model characteristics and runtime efficiency			
Metric	Baseline (FP32)	Quantized (INT8)	Distilled (FP32)
<i>Model Characteristics</i>			
Parameters (M)	270.4	270.4	21.5
Model Size (MB)	1031.4	127.9	82.0
Compression Ratio	1×	8.1×	12.6×
<i>GPU Performance</i>			
Throughput (samples/s)	31.9	—*	73.0
Latency (ms)	31.3	—*	13.7
Memory (MB)	8400.9	—*	295.8
<i>CPU Performance</i>			
Throughput (samples/s)	4.1	5.7	19.7
Latency (ms)	245.8	175.5	50.7
Memory (MB)	2695.9	2645.7	1282.4

(b) Detection performance and ablation studies					
Model	Variant	Acc.	P	R	F1
<i>Main Results</i>					
Baseline	Full Model	98.8	97.8	93.0	95.2
Quantized	Full Model	93.3	85.2	91.7	88.0
Distilled	Full Model	88.2	88.0	87.2	86.5
<i>Modality Ablation</i>					
Baseline	w/o Visual	96.8	93.9	94.8	94.3
	w/o Syntactic	97.0	94.6	94.7	94.6
	Text Only	96.7	94.0	94.6	93.8
<i>Compression Strategy Ablation</i>					
Quantization	Per-Tensor	92.3	83.6	91.8	86.9
	Per-Channel	90.5	79.9	83.0	81.3
Distillation	w/o Synergistic	79.3	74.8	84.3	78.5

*INT8 model optimized for CPU deployment only (fbgmm backend). Best results in each section are highlighted in **bold**.

**Fig. 3:** Per-class F1-scores across three clickbait categories.

gence by 9 percentage points, confirming that feature alignment losses are crucial for knowledge transfer.

Modality ablation on the baseline model reveals text dominance: removing visual or syntactic encoders causes minimal degradation, with text-only configuration maintaining 96.7% accuracy. This suggests multimodal fusion provides limited benefits for clickbait detection, contradicting the architectural complexity of current approaches.

Module-level sensitivity analysis, conducted by systematically quantizing individual layers, identifies critical vulnerability points. Middle Transformer layers (particularly Layer 6) show highest sensitivity to INT8 quantization, experiencing 1.33% accuracy drops for both FFN and attention

blocks. In contrast, embeddings and prediction heads demonstrate remarkable robustness, with the prediction head even showing slight improvement (+0.67%). This pattern suggests mixed-precision strategies should preserve middle-layer precision while aggressively quantizing peripheral components, potentially achieving better compression-accuracy trade-offs than uniform quantization.

5. CONCLUSION

We present a comprehensive framework for compressing multimodal clickbait detection models, achieving 8.1× compression through INT8 quantization (93.3% accuracy) and 92% parameter reduction via knowledge distillation (88.2% accuracy, 19.7 samples/s CPU). Our three-class Chinese dataset enables fine-grained analysis revealing that compression disproportionately impacts minority classes, with curiosity gap F1 dropping to 58.2% for distilled models. The baseline’s robust text-only performance (96.7%) and module-level sensitivity analysis identifying middle-layer vulnerability provide actionable insights for optimization. These results guide practical deployment: compression preserves overall accuracy but sacrifices minority class discrimination critical for detecting sophisticated manipulation. Cascaded architectures—lightweight screening followed by selective full-model processing—offer an immediate solution. Long-term, the field needs inherently efficient architectures rather than post-hoc compression, with our text-only results (96.7%) suggesting simpler models may suffice. This work provides benchmarks and strategies for deployable Chinese content moderation systems.

Acknowledgment.

This work was supported by the National Natural Science Foundation of China (Grant No.62406150), and the National Undergraduate Training Program for Innovation and Entrepreneurship.

6. REFERENCES

- [1] Mengxi Zhou, Wei Xu, Wenping Zhang, and Qiqi Jiang, “Leverage knowledge graph and GCN for fine-grained-level clickbait detection,” *World Wide Web*, vol. 25, no. 3, pp. 1243–1258, 2022.
- [2] Chuhan Wu, Fangzhao Wu, Tao Qi, and Yongfeng Huang, “Clickbait detection with style-aware title modeling and co-attention,” in *Proceedings of the 19th China National Conference on Chinese Computational Linguistics (CCL)*, 2020, vol. 12522, pp. 430–443.
- [3] Hei-Chia Wang, Martinus Maslim, and Hung-Yu Liu, “CA-CD: Context-aware clickbait detection using a new Chinese clickbait dataset with transfer learning methods,” *Data Technologies and Applications*, vol. 58, no. 2, pp. 243–266, 2024.
- [4] Xiaoyuan Yi, Jiarui Zhang, Wenhao Li, Xiting Wang, and Xing Xie, “Clickbait Detection via Contrastive Variational Modelling of Text and Label,” in *Proceedings of the 31st International Joint Conference on Artificial Intelligence (IJCAI)*, 2022, pp. 4475–4481.
- [5] Ye Wang, Yi Zhu, Yun Li, Liting Wei, Yunhao Yuan, and Jipeng Qiang, “Multi-modal soft prompt-tuning for Chinese clickbait detection,” *Neurocomputing*, vol. 614, pp. 128829, 2025.
- [6] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio, “Graph Attention Networks,” in *Proceedings of the 6th International Conference on Learning Representations (ICLR)*, 2018.
- [7] Benoit Jacob, Skirmantas Kligys, Bo Chen, Menglong Zhu, Matthew Tang, Andrew G. Howard, Hartwig Adam, and Dmitry Kalenichenko, “Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018, pp. 2704–2713.
- [8] Amir Gholami, Sehoon Kim, Zhen Dong, Zhewei Yao, Michael W. Mahoney, and Kurt Keutzer, “A survey of quantization methods for efficient neural network inference,” in *Low-power Computer Vision*, pp. 291–326, 2022.
- [9] Geoffrey Hinton, Oriol Vinyals, and Jeff Dean, “Distilling the knowledge in a neural network,” *arXiv preprint arXiv:1503.02531*, 2015.
- [10] Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf, “DistilBERT, a distilled version of bert: smaller, faster, cheaper and lighter,” *arXiv preprint arXiv:1910.01108*, 2019.
- [11] George Loewenstein, “The psychology of curiosity: A review and reinterpretation,” *Psychological Bulletin*, vol. 116, no. 1, pp. 75–98, 1994.
- [12] Celeste Kidd and Benjamin Y. Hayden, “The psychology and neuroscience of curiosity,” *Neuron*, vol. 88, no. 3, pp. 449–460, 2015.
- [13] Martin Potthast, Sebastian Köpsel, Benno Stein, and Matthias Hagen, “Clickbait Detection,” in *Proceedings of the 38th European Conference on IR Research (ECIR)*, 2016, vol. 9626, pp. 810–817.
- [14] Yiming Cui, Wanxiang Che, Ting Liu, Bing Qin, and Ziqing Yang, “Pre-training with whole word masking for Chinese BERT,” *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol. 29, pp. 3504–3514, 2021.
- [15] Wanxiang Che, Yunlong Feng, Libo Qin, and Ting Liu, “N-LTP: An Open-source Neural Language Technology Platform for Chinese,” in *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing (EMNLP): System Demonstrations*, 2021, pp. 42–49.
- [16] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever, “Learning Transferable Visual Models From Natural Language Supervision,” in *Proceedings of the 38th International Conference on Machine Learning (ICML)*, 2021, vol. 139, pp. 8748–8763.
- [17] Daya Khudia, Jianyu Huang, Protonu Basu, Summer Deng, Haixin Liu, Jongsoo Park, and Mikhail Smelyanskiy, “Fbgemm: Enabling high-performance low-precision deep learning inference,” *arXiv preprint arXiv:2101.05615*, 2021.

2027届本科生公能素质评价加分材料

	类别	分数（自填）	审核结果
思想成长与身心发展	理想信念坚定	5	
	热心服务集体	0	
	先锋作用突出	0.25	
志愿服务与社会实践	志愿服务	0.6	
	社会实践	0.5	
艺体素质与技能特长	代表学院参与文体竞赛	0.25	
	文体竞赛获奖	0	
到国际组织挂职实习	---	0	
总计	---	6.6	

材料统计截止时间：

学号： 2312220

2021级 2024年6月30日 ☐

姓名： 宋昊谦

2022级 2025年6月29日 ☐

2023级 2026年7月05日 ☐

所在 密码科学与技术班
班级：

(一) 思想成长与身心发展（上限10分）

1、理想信念坚定（上限5分）

在校期间思想表现情况自述：

在校期间，我始终注重思想政治素养和个人品德修养，认真学习党的理论和方针政策，自觉树立正确的世界观、人生观和价值观，具有较强的集体意识和责任意识。在日常学习和生活中，我尊敬师长、团结同学，诚实守信，积极参加学校和学院组织的各项活动，能够主动关心和帮助身边同学。

在学习方面，我始终保持认真严谨、勤奋踏实的态度，努力学习专业知识，不断提升自身专业能力和科研素养。在校期间积极参与科研实践，注重培养独立思考、分析问题和解决问题的能力，对科研工作具有较浓厚的兴趣和热情，希望今后能够继续深入学习和研究，在相关专业领域不断提升自己。

在纪律和生活方面，我严格遵守国家法律法规以及学校各项规章制度，自觉维护良好的学习和生活秩序，无任何违法违纪行为，未受到过任何处分。同时注重身心健康和全面发展，能够正确面对学习和生活中的压力，保持积极向上的心态。

总体而言，在校期间我在思想品德、学习科研、纪律意识和身心发展等方面均能够严格要求自己，具有较强的责任感和进取心，综合表现良好。

本人签字：

日期：

辅导员评分：

辅导员签字：

日期：

有无违反法律法规： 无

有无违反学校纪律处分规定： 无

有无发生思想品德问题并造成不良影响： 无

(一) 思想成长与身心发展（上限10分）

3、先锋作用突出



个人/团体奖项：	团体奖项	荣誉所属学院：	密码与网络空间安全学院
荣誉名称：	优秀班级团支部	获奖时间：	2025年6月
荣誉所属等级：	院级		

(二) 志愿服务与社会实践（上限6分）

1、志愿服务

(1) 志愿服务时长证明

基本资料

已报名活动

已参加活动

校外活动申报

校外活动

问题反馈

退出登录

基本资料

学号：2312220

姓名：宋昊谦

身份证号：411302200411270316

志愿者编号：411302001109357608

出生年月：2004-11-27

学院：密码与网络空间安全学院

专业：密码科学与技术

年级：2023级

入学时间：2023-09-02

毕业时间：2027-08-01

政治面貌：共青团员

手机号：16629057198

服务时长：25.5 小时

校外服务时长：0 小时

编辑

前三学年内志愿服务总时长（小时）： 25.5

除公能实践课要求外志愿服务总时长（小时）： 12

(二) 志愿服务与社会实践（上限6分）

- 1、志愿服务
- (1) 除公能实践课课要求18学时外，剩余志愿服务时长证明（需在统计时间范围内且为志愿服务平台导出证明）



志愿服务时长（小时）：

3

志愿服务时长（小时）：

3

志愿服务时长（小时）：

3

(二) 志愿服务与社会实践（上限6分）

- 1、志愿服务
- (1) 除公能实践课课要求18学时外，剩余志愿服务时长证明（需在统计时间范围内且为志愿服务平台导出证明）



志愿服务时长（小时）：

志愿服务时长（小时）：

志愿服务时长（小时）：

3

(二) 志愿服务与社会实践（上限6分）

2、社会实践（上限为2分）

(1) 社会实践次数证明

基本资料

申报活动

我的活动

已报名活动

已参加活动

活动反馈

中期答辩活动

申报基地

基地管理

申报书屋

书屋管理

已参加活动

搜索

活动名称	活动时间	认定时长	发布人	操作
2025年“梦圆南开 心系母校”的团体实践	2024-12-01 08:00到 2025-02-28 08:00	23 小时	张颖	查看 打印
2024年“梦圆南开 心系母校”的团体实践	2023-12-01 08:00到 2024-02-28 08:00	20 小时	张颖	查看 打印

大一寒假是否参加社会实践：是

大二寒假是否参加社会实践：是

大三寒假是否参加社会实践：否

大一暑假是否参加社会实践：否

大二暑假是否参加社会实践：否

大三暑假是否参加社会实践：否

(二) 志愿服务与社会实践（上限6分）

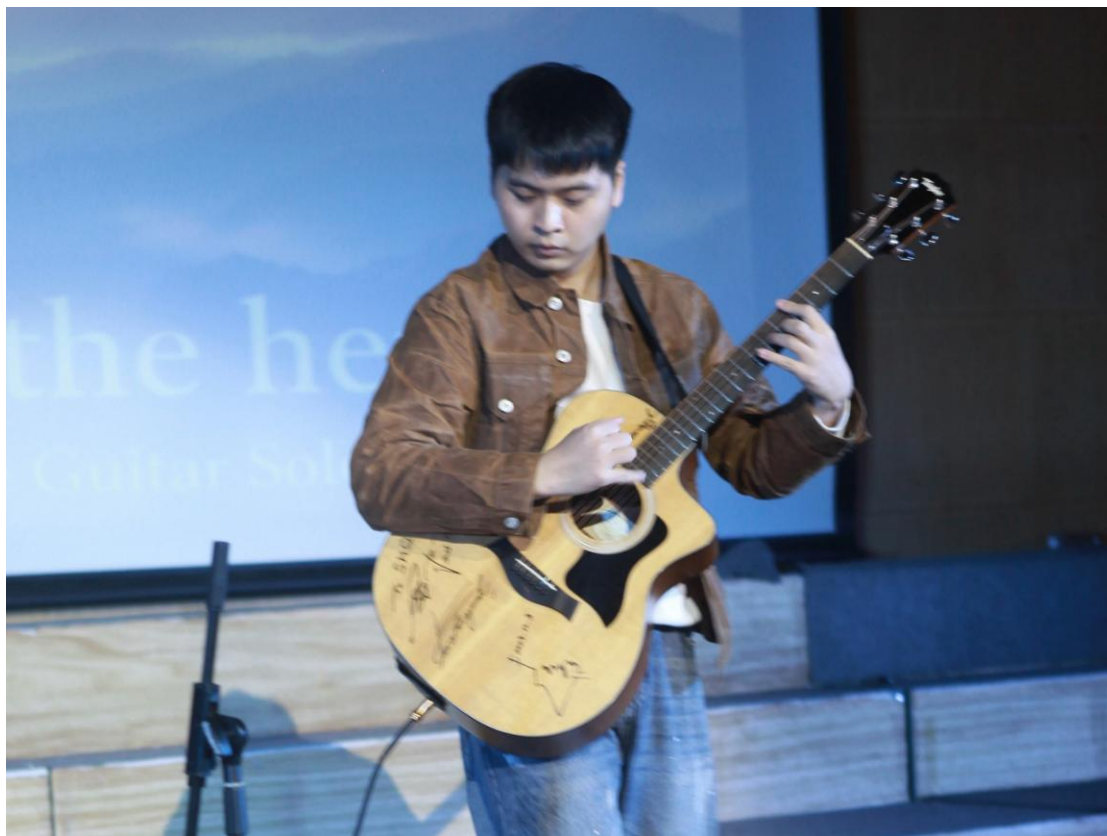
2、社会实践（上限为2分）

(1) 社会实践证明



(三) 艺体素质与技能特长（上限2分）

1、代表学校、学院参加艺术体育活动



参加活动名称： 计网学院2025年迎新演出

参加时间： 2025年11月9日

所属学院： 计算机学院和密码与网络空间安全学院

证明人： 吴洪柏