

ICPR 2026 Competition on Low-Resolution License Plate Recognition

Organized by: Rayson Laroca^{2,1}, Valfride Nascimento¹, and David Menotti¹

¹Federal University of Paraná, Curitiba, Brazil

²Pontifical Catholic University of Paraná, Curitiba, Brazil



PUCPR
GRUPO MARISTA

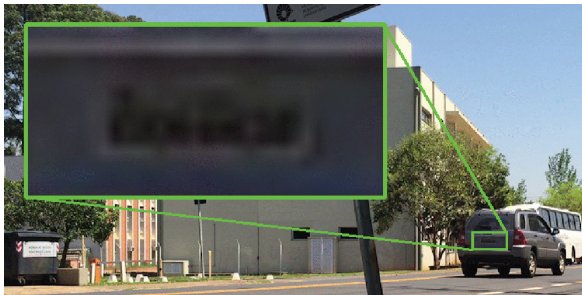
Introduction – ALPR on High-Quality Images

- Automatic License Plate Recognition (ALPR) performance has approached saturation on **high-quality images**.



Introduction – Real-World Degradations

- **However** → In real-world surveillance, images are frequently **degraded**:
 - Captured at **low resolution** due to hardware constraints or large distances.
 - Highly affected by **strong compression artifacts**, blur, and distortion.



Introduction – The LRLPR Challenge

- Despite its forensic importance, Low-Resolution LPR (LRLPR) remains a major challenge, with state-of-the-art recognition rates still around **50–60%**.

Introduction – The Gap in the Literature

- Most existing LRLPR studies rely on **synthetically degraded images** (e.g., bicubic downsampling from high-resolution samples).



Representative high-resolution (HR) and synthetic low-resolution (LR) image pairs.

Introduction – The Gap in the Literature

- **The Problem** → Synthetic models fail to capture the complex artifacts present in real operational scenarios.

Introduction – The Gap in the Literature

- **The Problem** → Synthetic models fail to capture the complex artifacts present in real operational scenarios.
- **Proposal:** Organizing the 1st Competition on LRLPR at ICPR 2026 to foster progress using **genuinely low-quality data** collected under real-world conditions.
 - The competition **website**: <https://icpr26lrlpr.github.io/>.

The 2026 LRLPR Competition

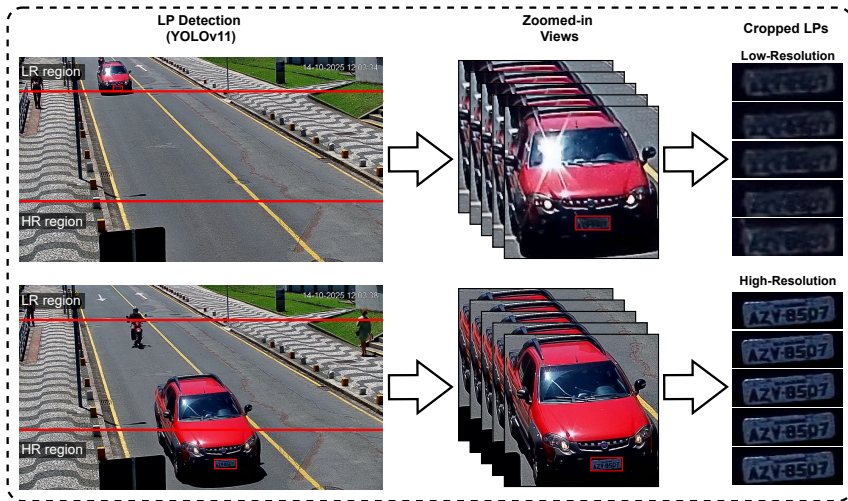
- Based on the **LRLPR Dataset**:
 - The **largest public dataset** featuring real low- and high-resolution LPs acquired from the same vehicles.
 - **20,000 training tracks** (5 LR + 5 HR images each).
 - **3,000 test tracks** (5 LR images each).

Dataset – Environmental Scenarios

- **Scenario A:** Collected under relatively controlled conditions (daylight, no rain).
- **Scenario B:** Collected for this competition, covering a **broader range of conditions** (rain, nighttime). Encompasses a different camera angle.



Dataset – Construction Pipeline



- Farthest frames → **LR samples**; Nearest frames → **HR samples**.

Dataset – LR vs. HR Tracks

- Each training track has **5 LR** and **5 HR images** of the same vehicle.

Track 1: **LR** Images



Track 1: **HR** Images



Track 2: **LR** Images



Track 2: **HR** Images



- **Two Phases via Codabench:**

- ① **Public Test Phase:** subset of 1,000 of the 3,000 test tracks with public leaderboard feedback.
- ② **Blind Test Phase:** full test set of 3,000 tracks with a private leaderboard and a maximum of 3 submissions to prevent probing.

- **Two Phases via Codabench:**

- ① **Public Test Phase:** subset of 1,000 of the 3,000 test tracks with public leaderboard feedback.
- ② **Blind Test Phase:** full test set of 3,000 tracks with a private leaderboard and a maximum of 3 submissions to prevent probing.

- **Evaluation Metrics:**

- **Primary: Recognition Rate** (requires an exact match between predicted and ground-truth strings).
- **Secondary: Confidence Gap** (breaks ties; measures the model's ability to differentiate between correct and incorrect predictions).
 - Confidence Gap = the difference between the mean confidence assigned to correct predictions and incorrect predictions.

- The competition attracted global interest:

Participation & Engagement

- The competition attracted global interest:
 - **269 teams** from **41 countries**.
 - Ranked among the **top 10 competitions on Codabench** by number of participants.

Participation & Engagement

- The competition attracted global interest:
 - **269 teams** from **41 countries**.
 - Ranked among the **top 10 competitions on Codabench** by number of participants.
- **Submissions:**
 - **118 teams** submitted valid entries in the Public Test Phase.
 - **99 teams** participated in the final Blind Test Phase.

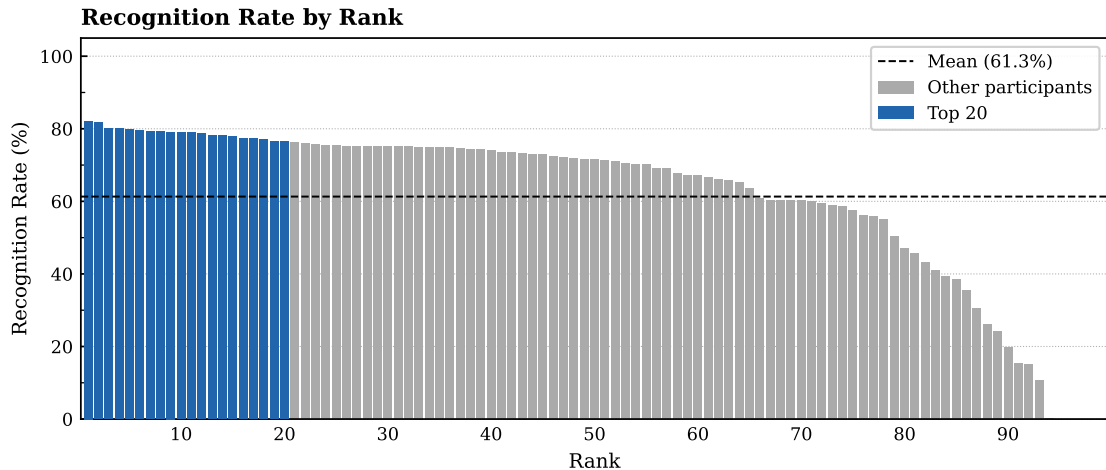
Results – Blind Test Phase

- The top of the leaderboard was **highly competitive**.

Rank / Team	Recognition Rate (↑)	Confidence Gap (↑)
1 st – DLmath	82.13%	6.67%
2 nd – AIO_JiangnamCoffee	81.73%	3.75%
3 rd – OpenOCR	80.17%	2.38%
4 th – CAP2	80.10%	14.86%
5 th – UIT-MeoBeo	79.83%	5.93%
Top 10 Gap	Separated by only 3.13 percentage points	

- Even the top-performing method failed on **17.87%** of test tracks under the strict exact-match protocol.

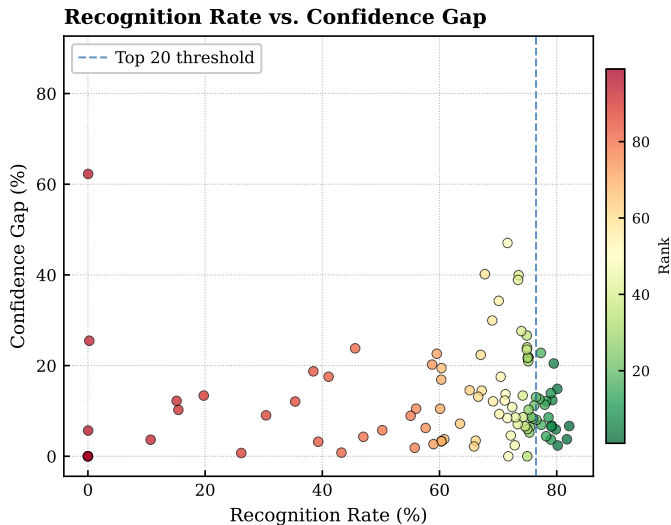
Results – Performance Distribution



Recognition Rate achieved by all teams. The top 20 teams are highlighted in blue.

Results – Confidence Gap

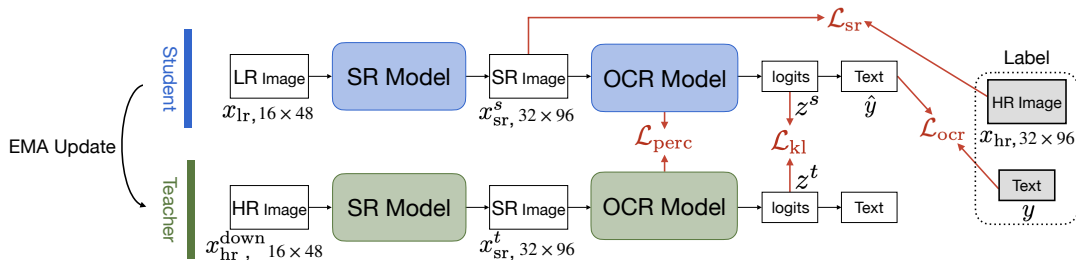
- Methods with similar Recognition Rates can differ **substantially** in Confidence Gap.



Top Approaches – 1st Place (DLmath)

Team: D. Kim, S. Chung, S. Bae, U. Seo, and S. Oh (Korea University)

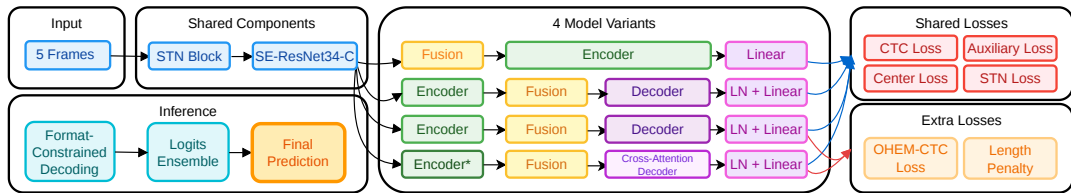
- **Teacher-student framework** jointly training a super-resolution model and an OCR model.
- **Student branch** takes LR inputs; **Teacher branch** receives downsampled HR images.
- **Inference:** Logits from all 5 LR frames in a track are summed before decoding.



Top Approaches – 2nd Place (AIO_JiangnamCoffee)

Team: C. M. Phung and M. G. Vo (UIT, HCMUT, and VNU, Vietnam)

- **Four-stage framework:** STN for alignment, SE-ResNet-C backbone, Transformer encoder for sequence modeling, and a CTC head.
- A CNN-based attention module estimates frame quality and fuses the 5 input frames into a weighted representation.
- Ensembled 4 model variants differing in multi-frame fusion strategy, Transformer configuration, and auxiliary losses.



Top Approaches – 3rd Place (OpenOCR)

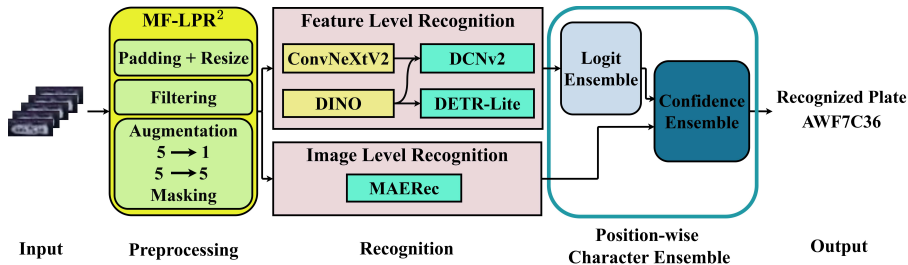
Team: X. Ye, Y. Du, Y. Su, and Z. Chen (Fudan Univ. & Shanghai Key Lab., China)

- Treated LRLPR directly as a **robust scene text recognition** problem without a full super-resolution stage.
- **Backbone:** SVTRv2-AR (an attention-based autoregressive model).
- **Multi-Frame Fusion:** The 5 per-frame predictions were aggregated using a **character-level voting strategy** combining prediction frequency and confidence scores.
- Ensembled 4 SVTRv2-AR models using different datasets and initialization strategies.

Top Approaches – 4th Place (CAP2)

Team: S. Heo, H. Lee, and K. Na (Handong Global University, South Korea)

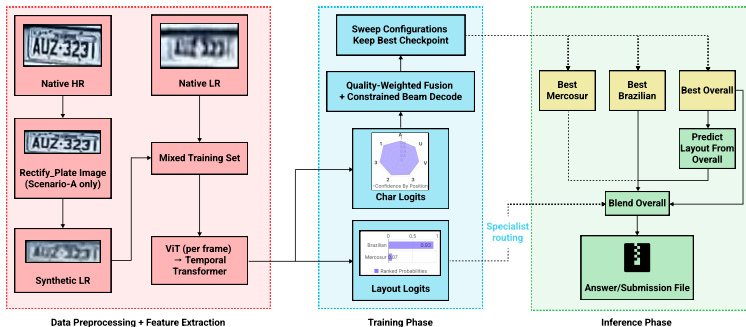
- **Multi-stage pipeline:** geometry-aware preprocessing, dual-stream recognition, and a position-wise ensemble.
- **Preprocessing:** Extends MF-LPR2, but treats each frame as an **independent restored candidate** instead of fusing them into one.
- **Dual-Stream Recognition:** Combines multiple feature-level extractors (ConvNeXtV2, DINOv2/v3) and image-level recognizers.



Top Approaches – 5th Place (UIT-MeoBeo)

Team: K. Nguyen, S. Pham, D. Phung, T. Le, and V. Tran (UIT & VNU, Vietnam)

- **Multi-stage, multi-frame OCR pipeline** using a PE-Core-L frame encoder and a two-layer temporal Transformer.
- **Dual Prediction Heads:** Jointly estimate the LP text and the layout.
- **Structure-Aware Decoding:** Fixed 7-character output with position-wise letter/digit constraints based on the estimated layout.



- **No single dominant design:**
 - Top teams adopted substantially different strategies (e.g., explicit super-resolution vs. direct recognition from LR).

- **No single dominant design:**
 - Top teams adopted substantially different strategies (e.g., explicit super-resolution vs. direct recognition from LR).
- **The Key to Success** → Effective use of the five-frame track structure.
 - Temporal fusion, voting, logit aggregation, and ensemble-based integration were critical across the strongest submissions.

- **No single dominant design:**
 - Top teams adopted substantially different strategies (e.g., explicit super-resolution vs. direct recognition from LR).
- **The Key to Success** → Effective use of the five-frame track structure.
 - Temporal fusion, voting, logit aggregation, and ensemble-based integration were critical across the strongest submissions.
- **Confidence Gap matters:** Improving systems requires not just accuracy, but making confidence estimates more informative for forensic workflows.

- We established the first large-scale international benchmark specifically focused on **LRLPR with real low-quality data**.

- We established the first large-scale international benchmark specifically focused on **LRLPR with real low-quality data**.
- The competition provided a snapshot of the methodological diversity driving progress, offering a clear basis for future baselines.

- We established the first large-scale international benchmark specifically focused on **LRLPR with real low-quality data**.
- The competition provided a snapshot of the methodological diversity driving progress, offering a clear basis for future baselines.
- **Future Directions:**
 - Improved multi-frame modeling.
 - Better confidence estimation.
 - Tighter integration between restoration and recognition.



Thank you!

<https://icpr26lrlpr.github.io/>

