

# SCIENCE CHINA Information Sciences

## Call for Papers

### Special Topic: Vision-Language-Action and World Models



In recent years, vision-language-action (VLA) and World Models have attracted widespread attention and undergone explosive growth in both academia and industry. VLA models introduce “action” as a native output modality into large multimodal models, while World Models learn the underlying dynamics of an environment to predict future states, imagine the consequences of actions, and plan through simulation. Together, these two paradigms have demonstrated substantial potential in robotic manipulation, navigation, autonomous driving, embodied interaction, and interactive simulation.

This special topic aims to encourage and publish cutting-edge original research, surveys, and insightful benchmark studies on VLA models and World Models. Topics of interest include, but are not limited to:

(1) **Unified Representation and Learning for VLA:** Research on cross-modal alignment, unified modeling, and end-to-end learning mechanisms across visual perception, language understanding, and action generation.

(2) **Construction and Learning of World Models:** Research on modeling methods for environment states, object relations, physical laws, and dynamic evolution processes, to enhance an agent's understanding and prediction of the external world.

(3) **Fusion Architectures of VLA and World Models:** Exploration of the deep integration between the predictive and forward-simulation capabilities of World Models and the perception, reasoning, and control capabilities of VLA models, toward closed-loop embodied intelligence systems.

(4) **Long-Horizon Task Planning and Execution:** Research on task decomposition, hierarchical planning, subgoal generation, execution monitoring, and dynamic replanning methods for complex, long-horizon tasks.

(5) **Memory and Continual Learning for VLA and World Models:** Research on modeling episodic, event, semantic, and skill memory to support knowledge accumulation across tasks, scenes, and time scales.

(6) **Generative World Models and Future Prediction:** Research on scene generation, video prediction, state transition modeling, trajectory generation, and simulation of action consequences, providing simulatable environments for agent decision-making and planning.

(7) **Spatial Intelligence, Physical Reasoning, and Causal Modeling:** Research on 3D spatial understanding, affordance analysis, modeling of physical laws, and causal reasoning, to enhance an agent's understanding of interaction objects and action outcomes.

(8) **Embodied AI Data, Simulation, and Training Methods:** Research on key methods including multimodal data collection, automatic annotation, synthetic data generation, simulation-based training, reinforcement learning, and self-supervised learning.

(9) **Efficient, Trustworthy, and Safe Techniques for VLA and World Models:** Research on parameter-efficient training, model compression, inference acceleration, robustness, interpretability, uncertainty estimation, and safety constraint mechanisms.

(10) **Evaluation and Industrial Applications of VLA and World Models:** Research on unified benchmarks, evaluation metrics, and system validation methods, as well as applications in robotics, autonomous driving, intelligent manufacturing, logistics, and service scenarios.

## Submission

The papers should be prepared using the SCIS template, and should be submitted online through the manuscript submission system of the SCIENCE CHINA Information Sciences. The submission website is

<https://mc03.manuscriptcentral.com/scis>.

You should choose **Special Topic: Vision-Language-Action and World Models**. Information and guidelines on preparation of manuscripts are available on the journal website: <http://scis.scichina.com>.

## Important Dates

Submission deadline: **November 30, 2026**

First Round of Reviews: January 31, 2027

Final Acceptance: March 31, 2027

Publication: May 1, 2027

## Guest Editors

**Xiang BAI (白翔)**, Huazhong University of Science and Technology, China

**Weinan ZHANG (张伟楠)**, Shanghai AI Lab and SJTU, China

**Wei SHEN (沈为)**, Shanghai Jiao Tong University, China

**Hengshuang ZHAO (赵恒爽)**, The University of Hong Kong, China

**Dingkang LIANG (梁定康)**, Huazhong University of Science and Technology, China

## Contact

**Kai JIANG (蒋恺)**, Scientific Editor,  
SCIENCE CHINA Information Sciences Editorial Office,  
[jiangkai@scichina.com](mailto:jiangkai@scichina.com), 010-64015683