

Testing and Evaluation of Autonomous Driving Systems

From Simulated to Real-world Test Environments

Andrea Stocco

SIESTA summer school

02.09.2026



whoami

Cyber Security at *Tallinn University of Technology* and *University of Tartu*, Estonia



PhD in Computer Science at *University of Genova*, Italy



PostDoctoral Researcher of the *Software Analysis and Testing* group at University of British Columbia, Canada



PostDoctoral Researcher of the *Software Institute* at Università della Svizzera italiana, Switzerland



whoami

fortiss



*Lead of the Automated Software
Testing (AST) Field of Competence*

*Chair of Software Engineering for
Data-intensive Applications*

Email

stocco@fortiss.org
andrea.stocco@tum.de

Website/X/LinkedIn

tsigalko18.github.io/

 @tsigalko18

[www.linkedin.com/in/
andreastocco/](https://www.linkedin.com/in/andreastocco/)

Acknowledgments



**Stefano Carlo
Lambertenghi**



**Xingcheng
Chen**



**Lev
Sorokin**



**Oliver
Weiß**



**Davide Yi
Xian Hu**



**Manfredi
Napolitano**

<https://github.com/ast-fortiss-tum/>

And many other colleagues and co-authors...

Automated Software Testing

Main domain areas

- **Traditional Software Systems**

- Search-based Test Generation
- Regression Testing (e.g., test prioritization, test minimization)

- **Web Applications**

- End-to-End GUI Testing
- Automated Crawling

- **Deep Learning Systems**

- Test Generation
- Failure Prediction (black-box, white-box)

- **Cyber Physical Systems**

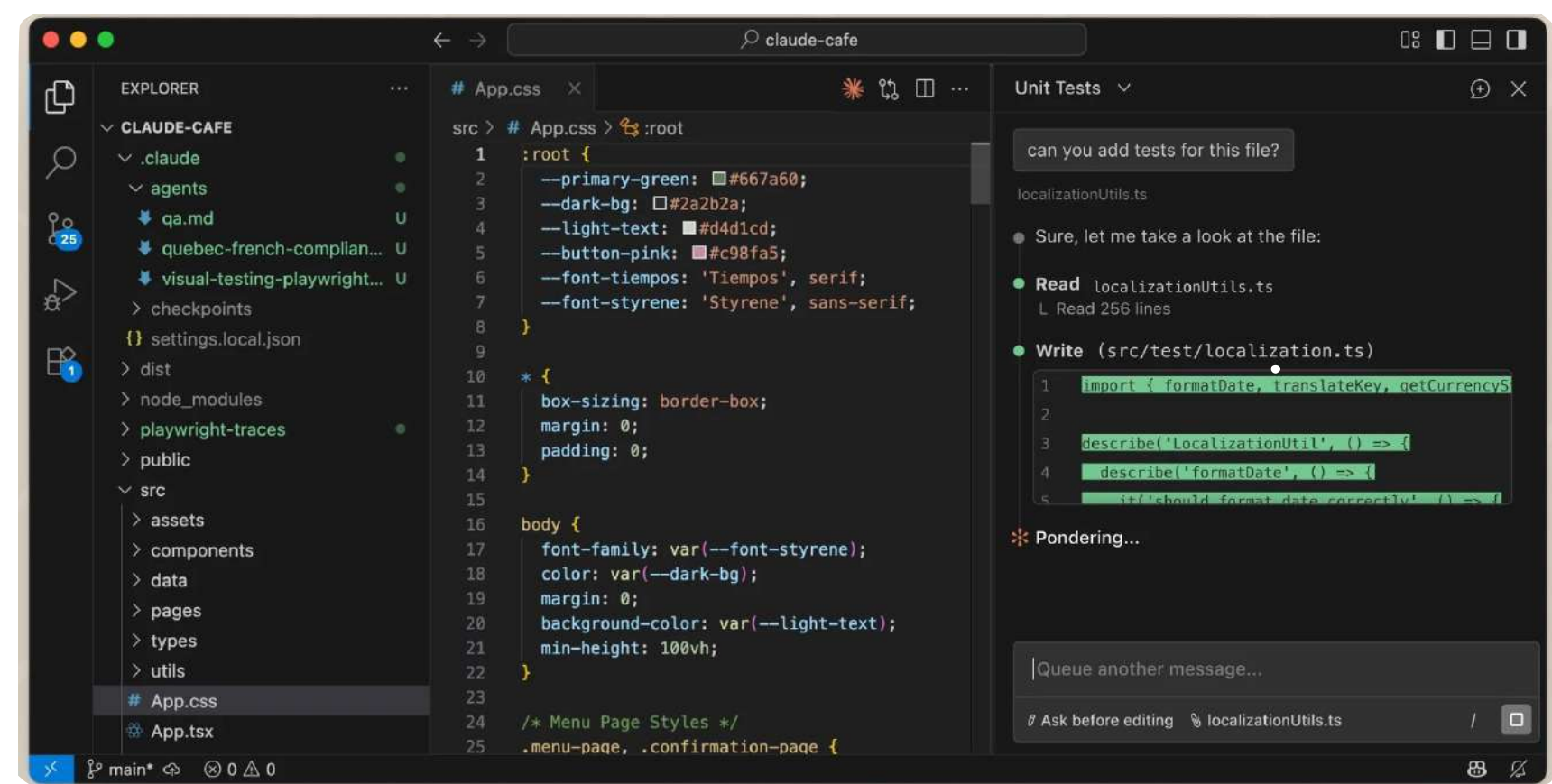
- Reality Gap Assessment and Mitigation
- Automotive, UAVs, Elevators

- **LLM Systems**

- Test Generation
- Conversational Systems, UAVs, Elevators



AI-based Systems



1982: What if we could build an intelligent car?



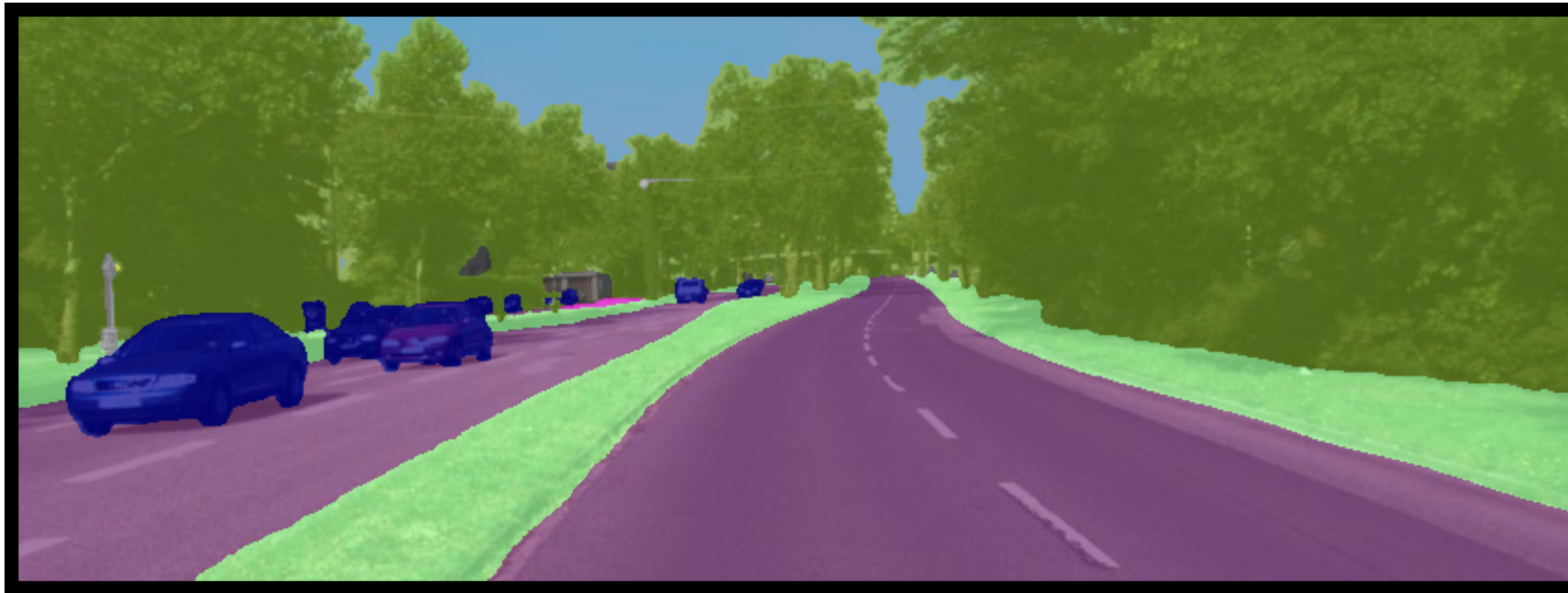
2026: We can. Now, how can we trust it?



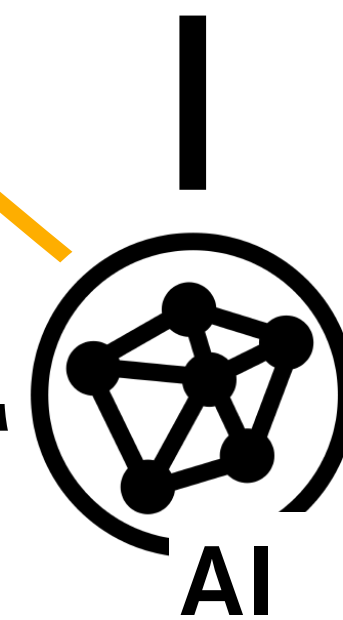
Camera/LiDAR/Radar



Deep Neural Networks
Visual Language Models



MODULAR | Segmentation



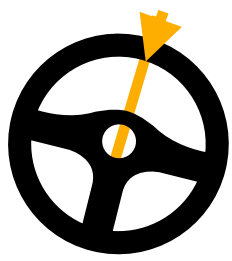


MODULAR
| Segmentation

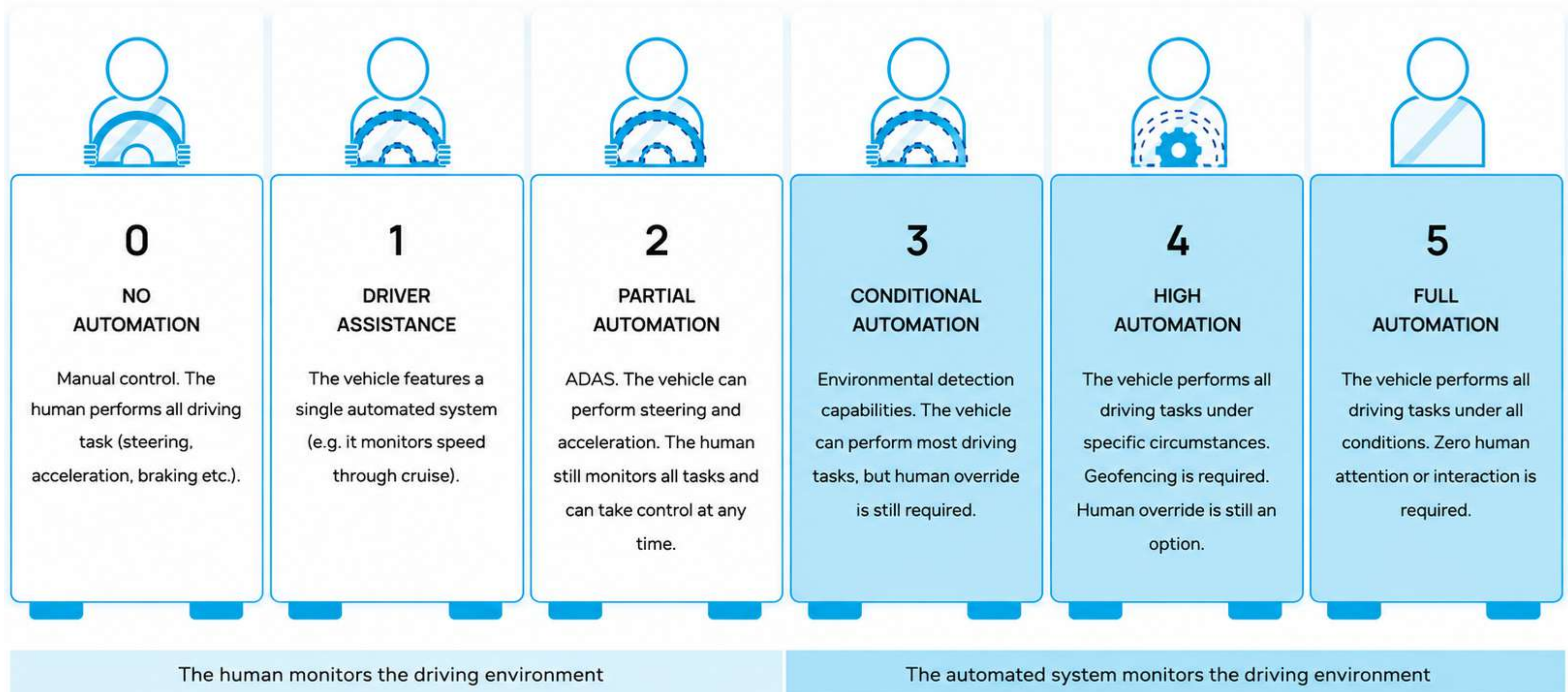
END-2-END

Lane-keeping +
Adaptive Cruise Control



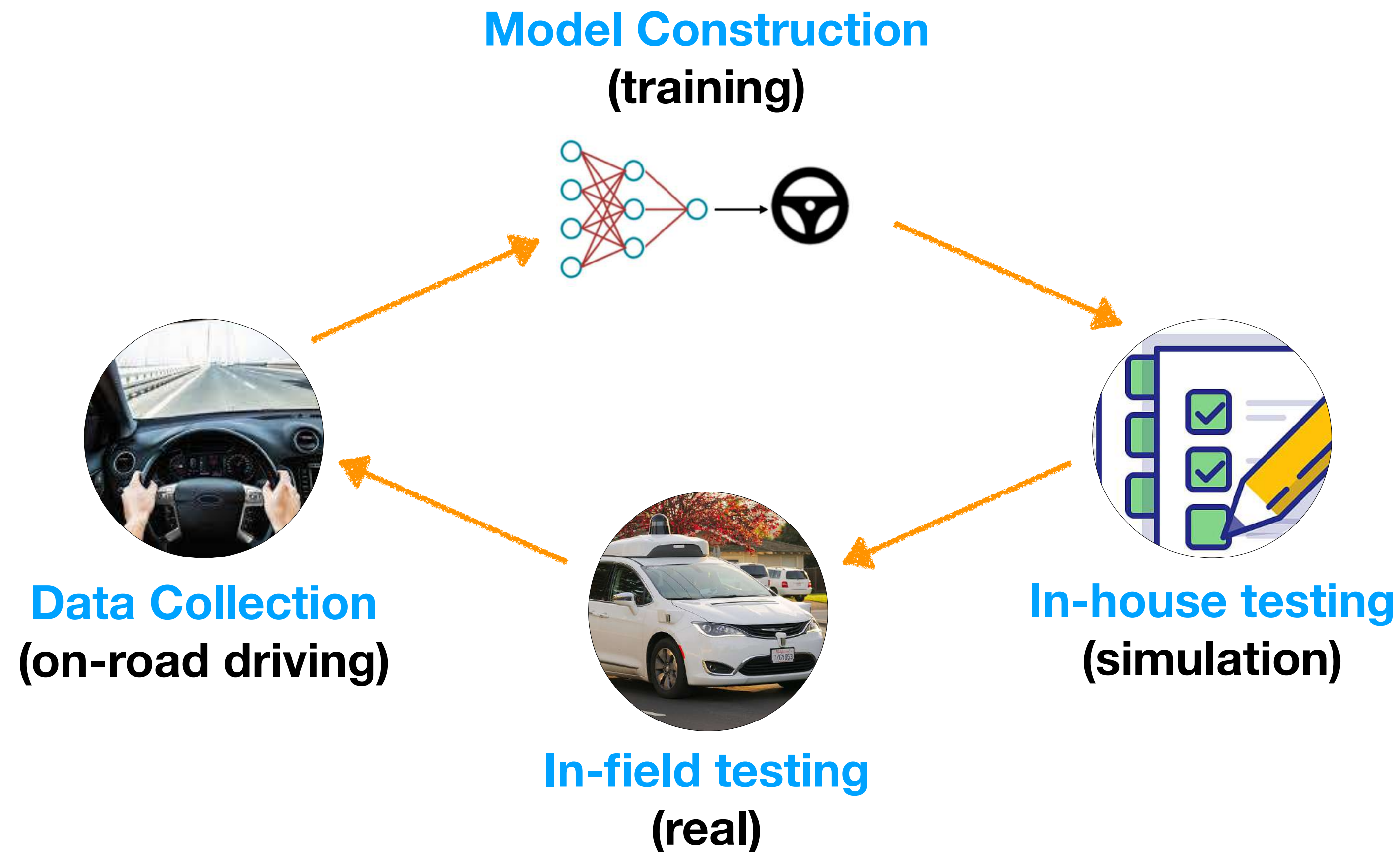
Steer 0.1 
Throttle 0.3 

LEVEL OF DRIVING AUTOMATION



Automated Driving System (ADS) testing

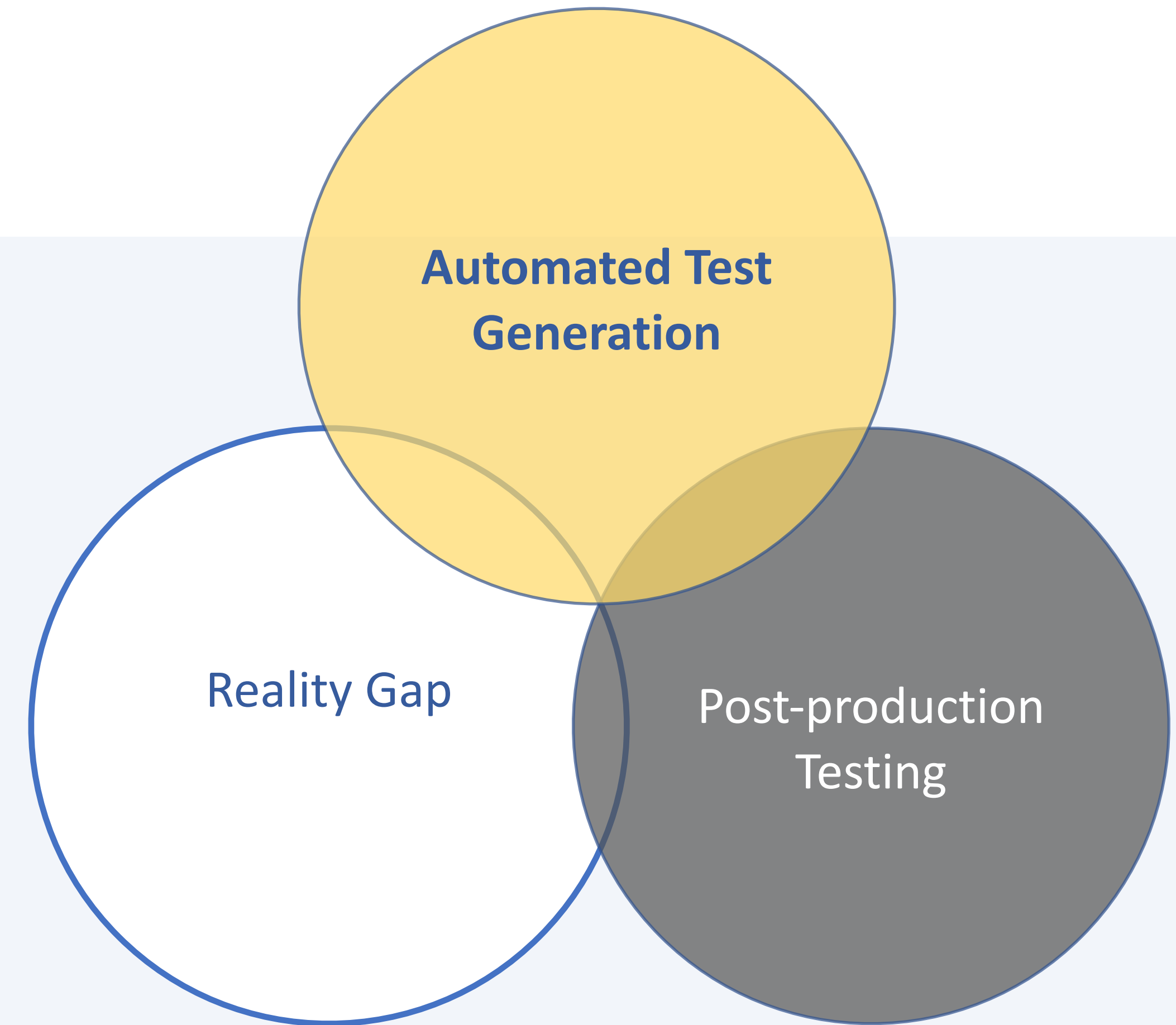
How to ensure that an ADS system is ready for deployment?



ADS Testing

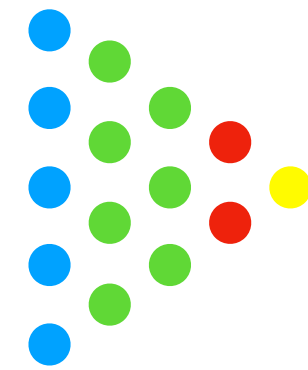
Main research topics

- **Automated Test Generation**
How can we automatically generated complex scenario-based tests efficiently and effectively?
- **Reality Gap**
How can we measure and mitigate the reality gap between simulation and in-field testing?
- **Post-production Testing**
How to ensure a high dependability of deep neural network driven-cyber-physical systems (CPS) in production?

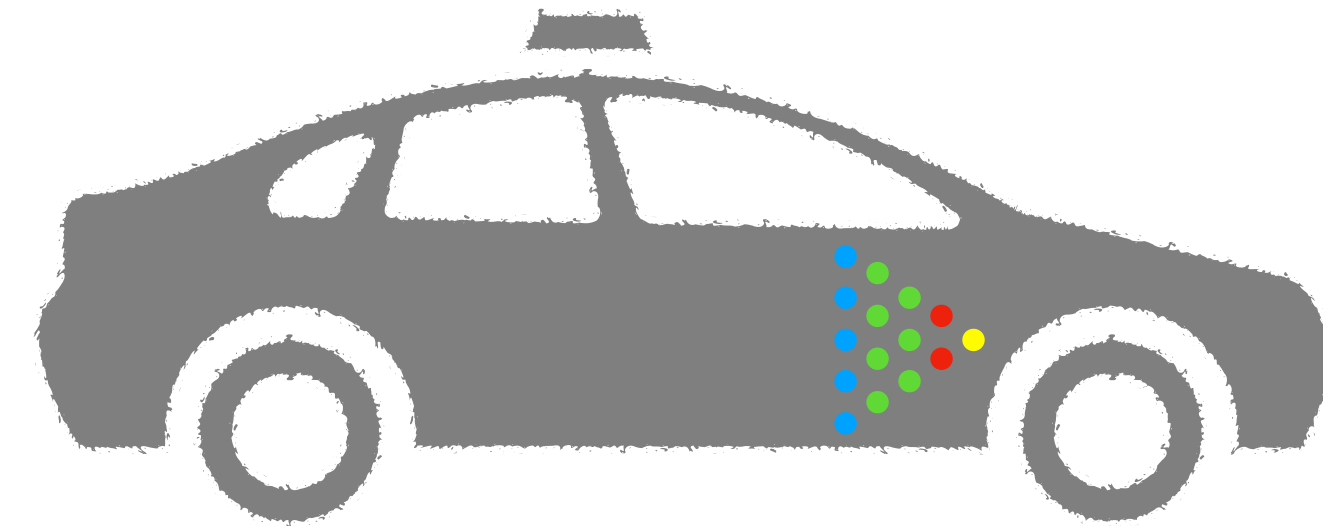


Model vs System level Testing

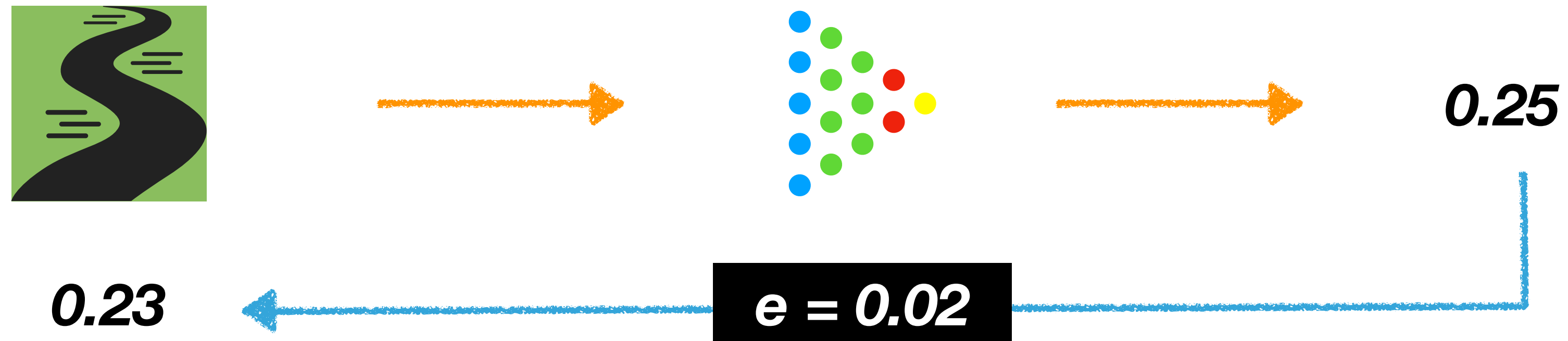
Model



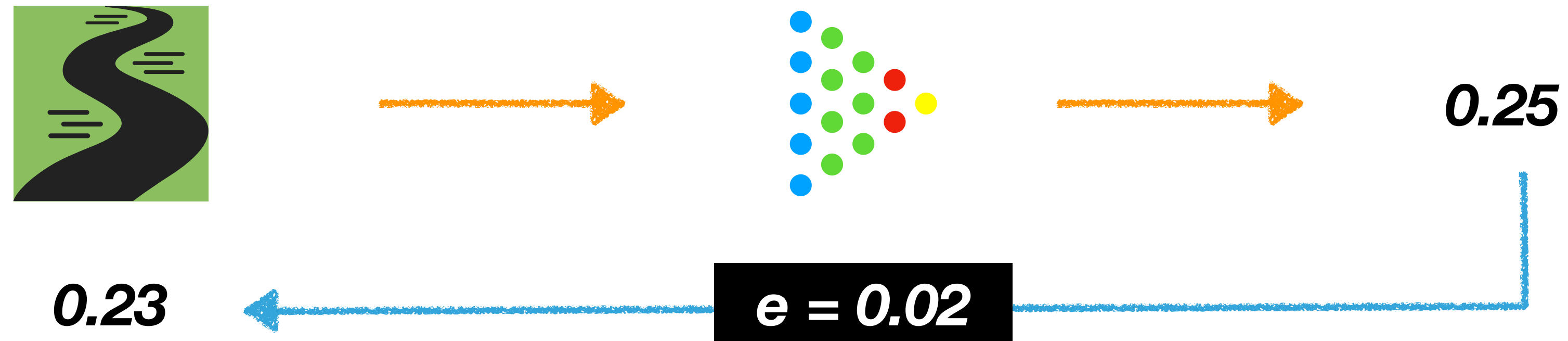
System



Model level Testing



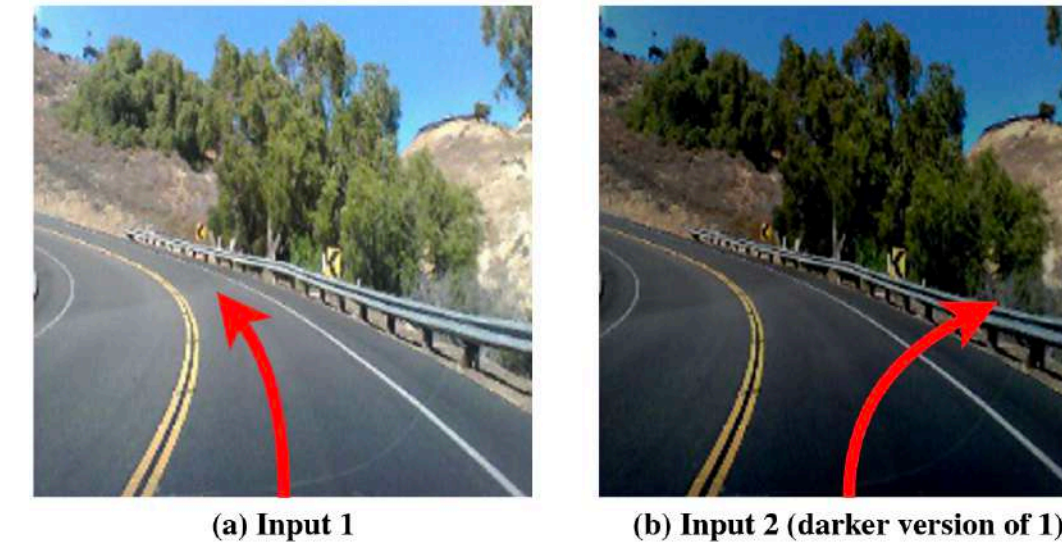
Model level Testing



- Quick Feedback
- Cheap
- No driving
- Simple Oracle

Model level Testing Literature

DeepXplore
Pei et al., 2017



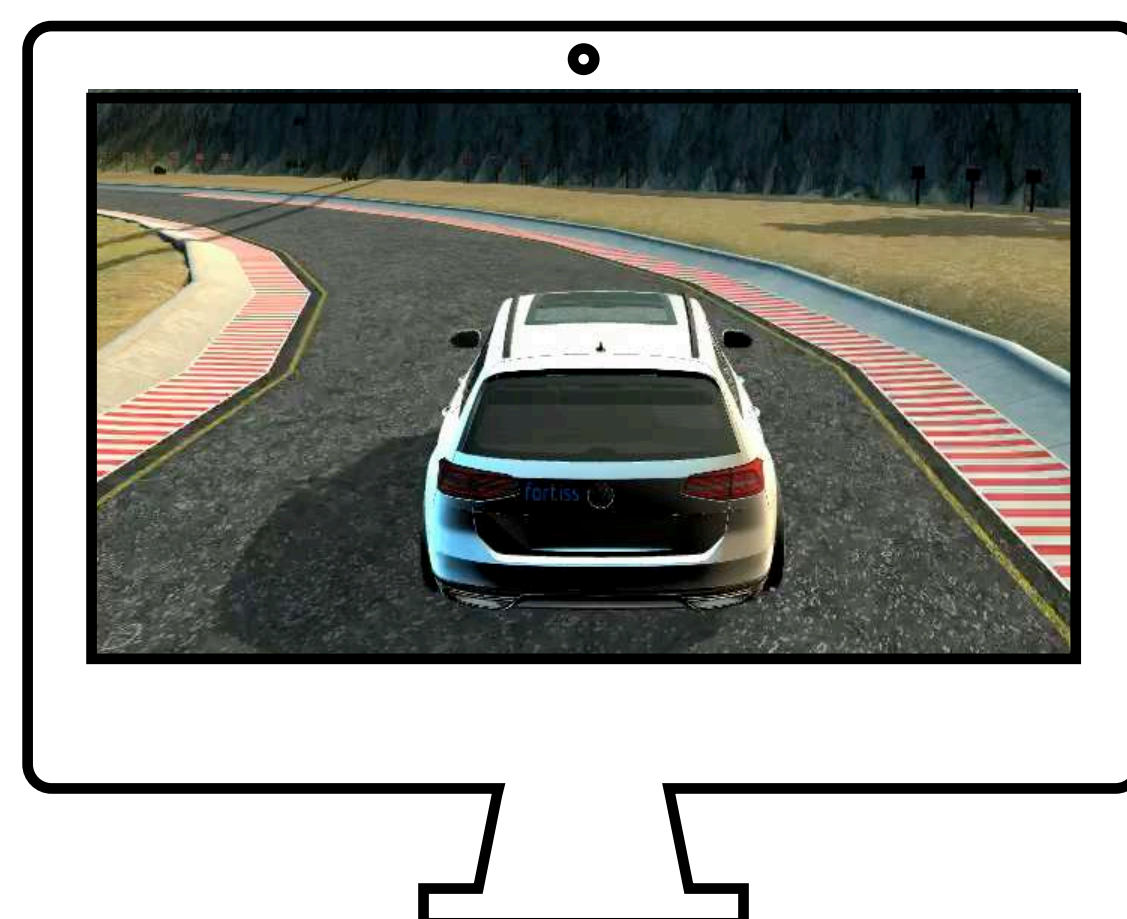
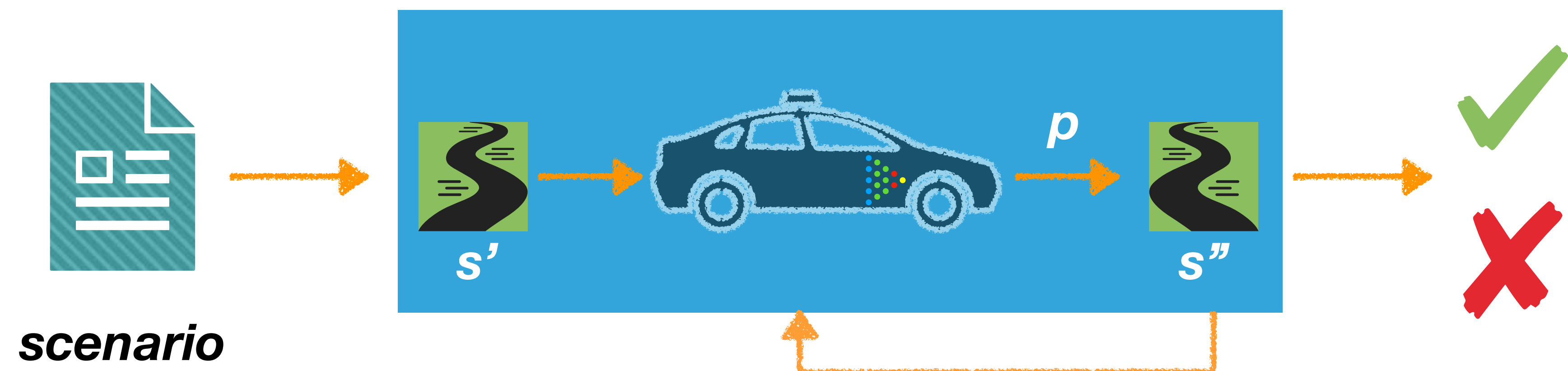
DeepTest
Tian et al., 2017



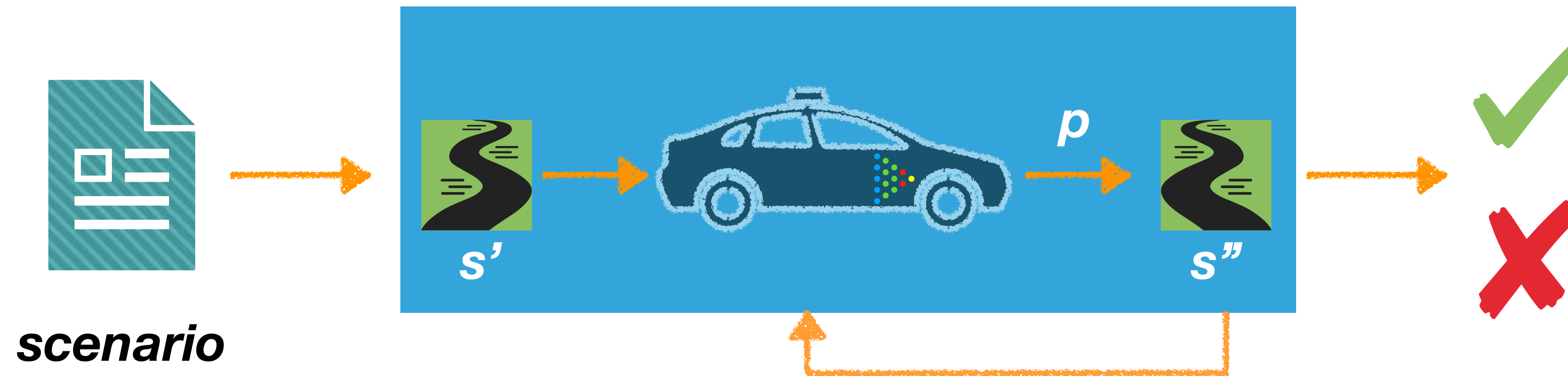
DeepRoad
Zhang et al., 2018



System level Testing



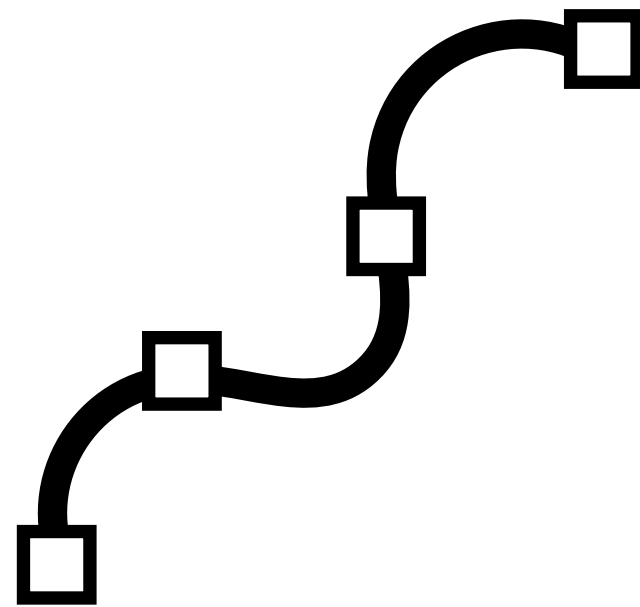
System level Testing



- Driving
- More realistic

- Time consuming
- Complex Oracle

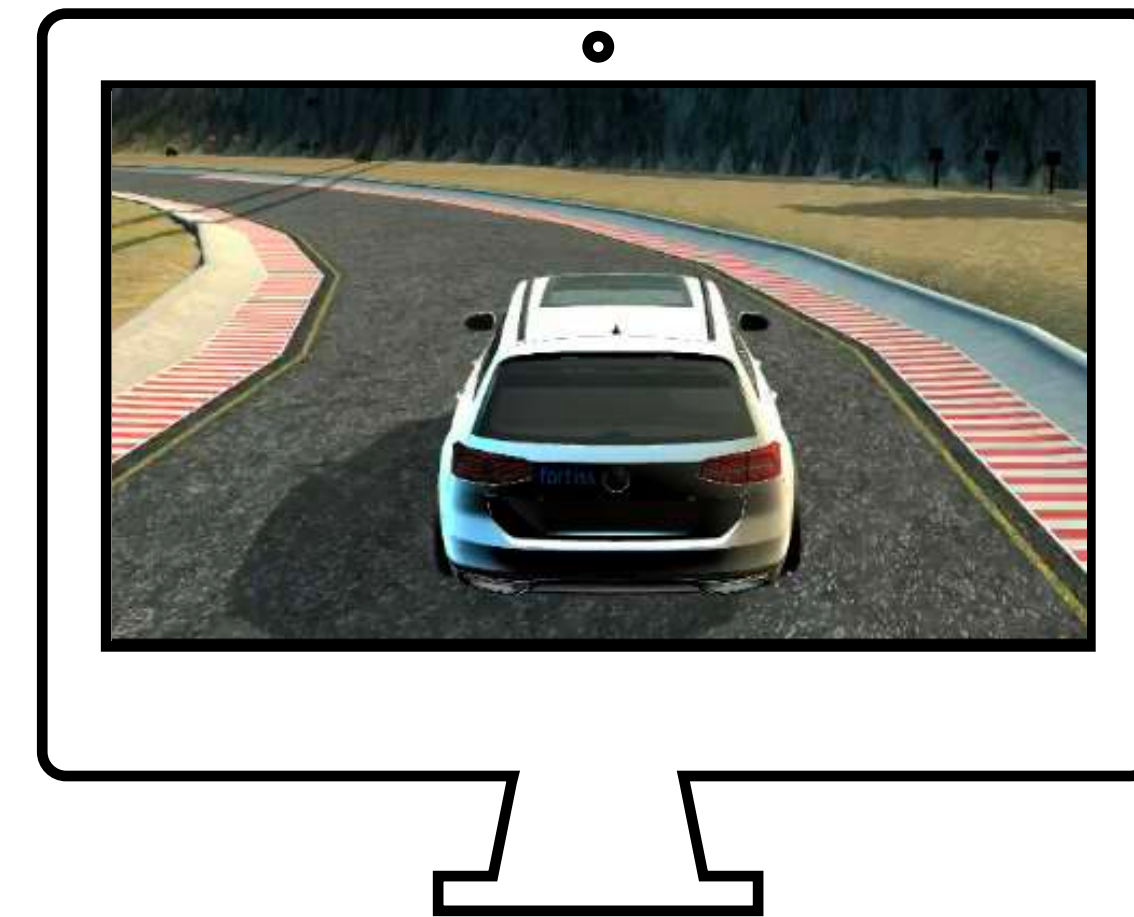
Test Scenarios for ADS



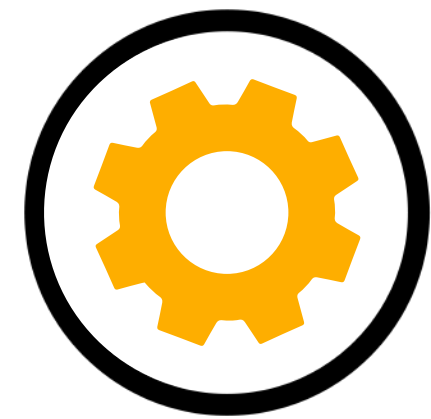
Lane-keeping Driving Assistance System



Modular or End-to-End Driving System



Testing for Robustness



PerturbationDrive



**33 perturbations
5 intensities**



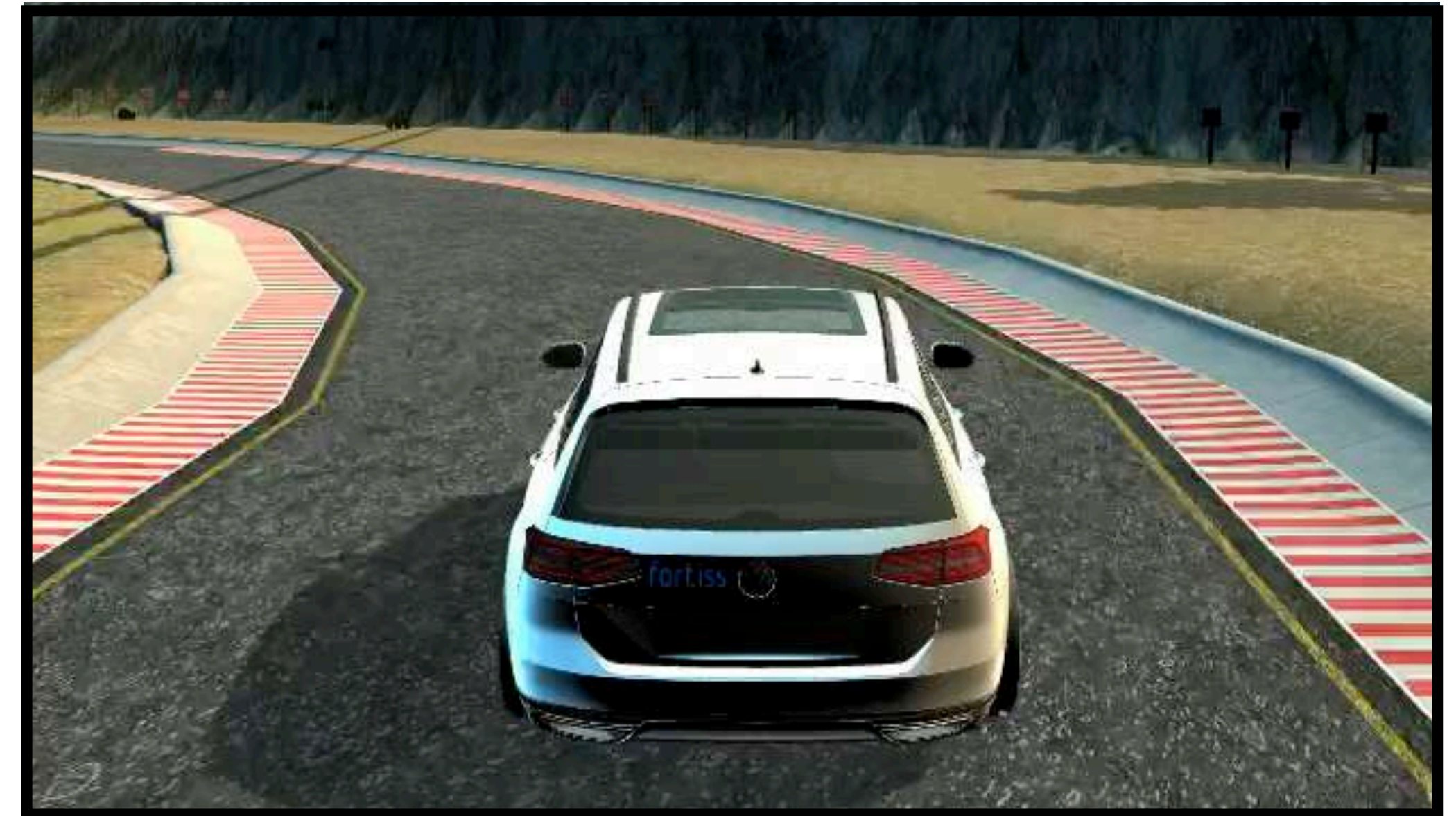
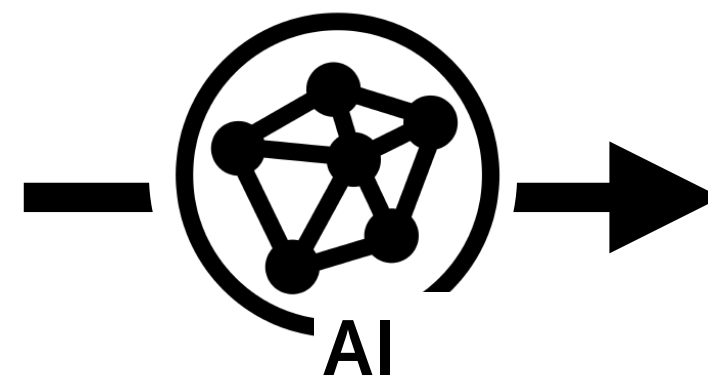
**Model & system
level**



**Multiple
simulators**

```
> apply("Phase scrambling",intensity=4)
```

PerturbationDrive





SIM 1
Small DT

Phase scrambling



↓ 54%
Success
rate

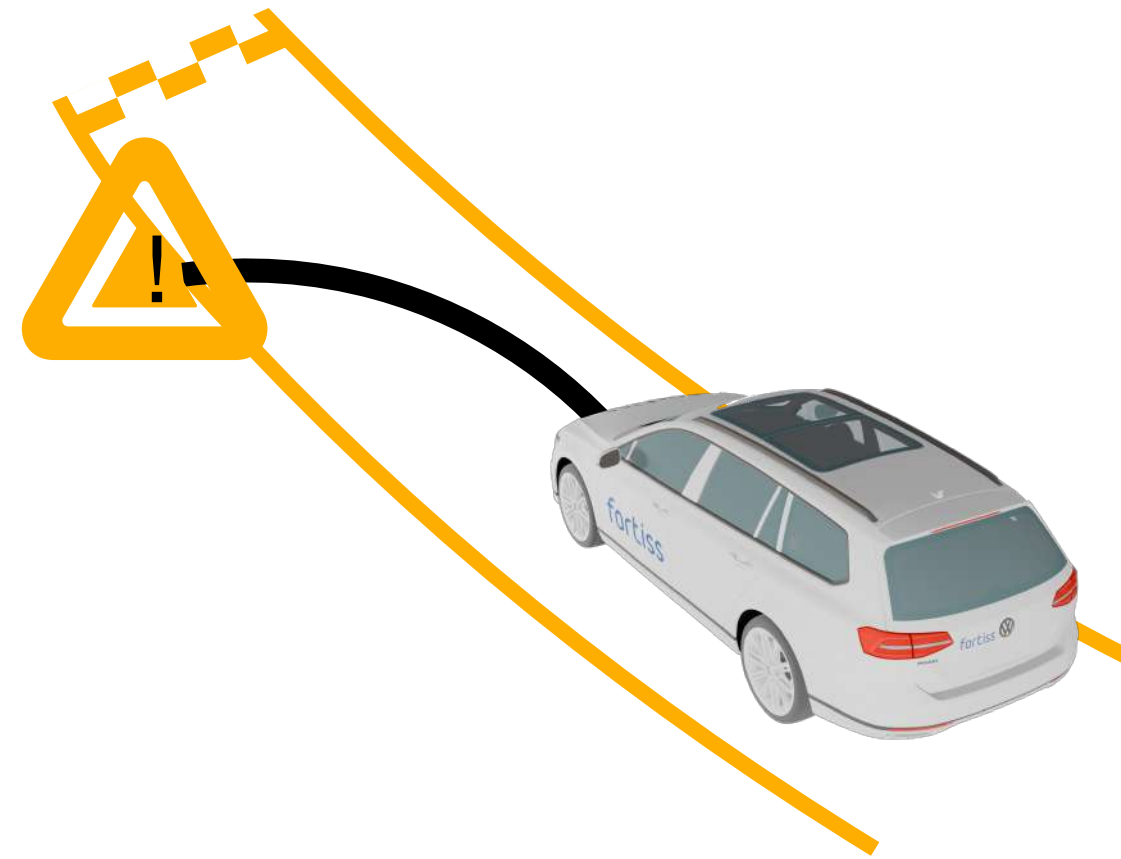
SIM 2
Udacity

False color

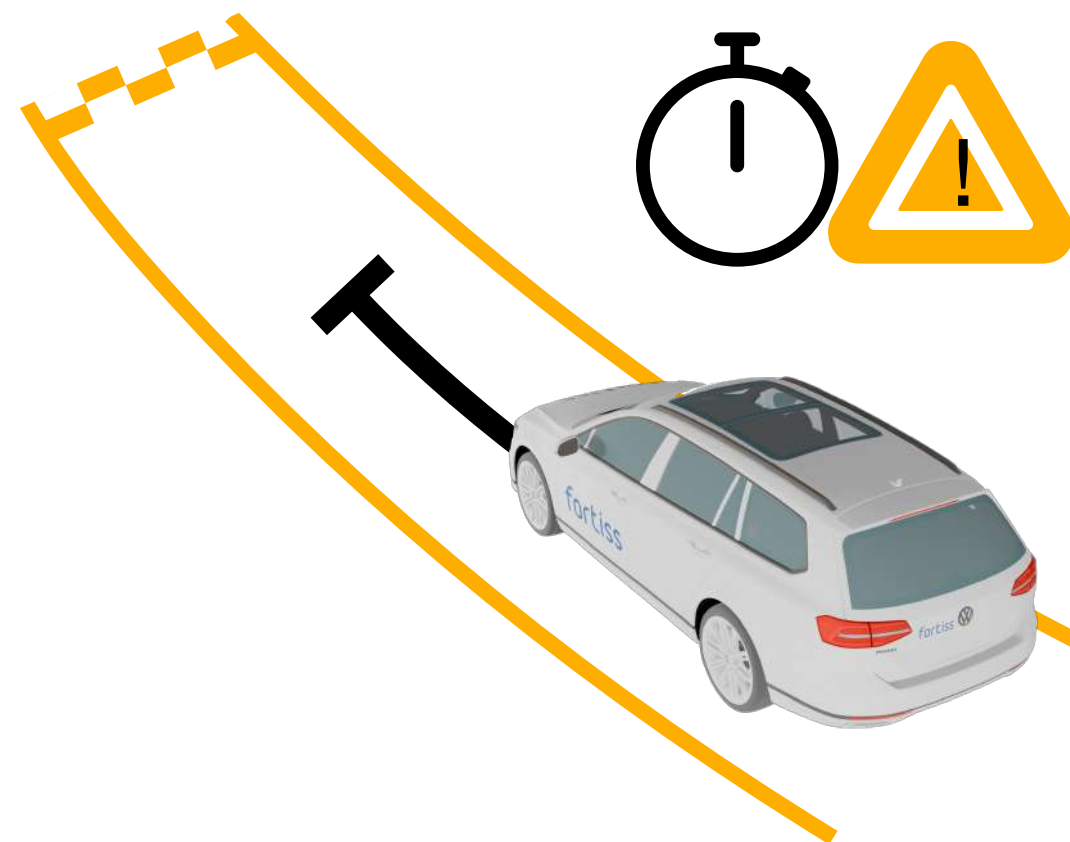


↓ 76%
Success
rate

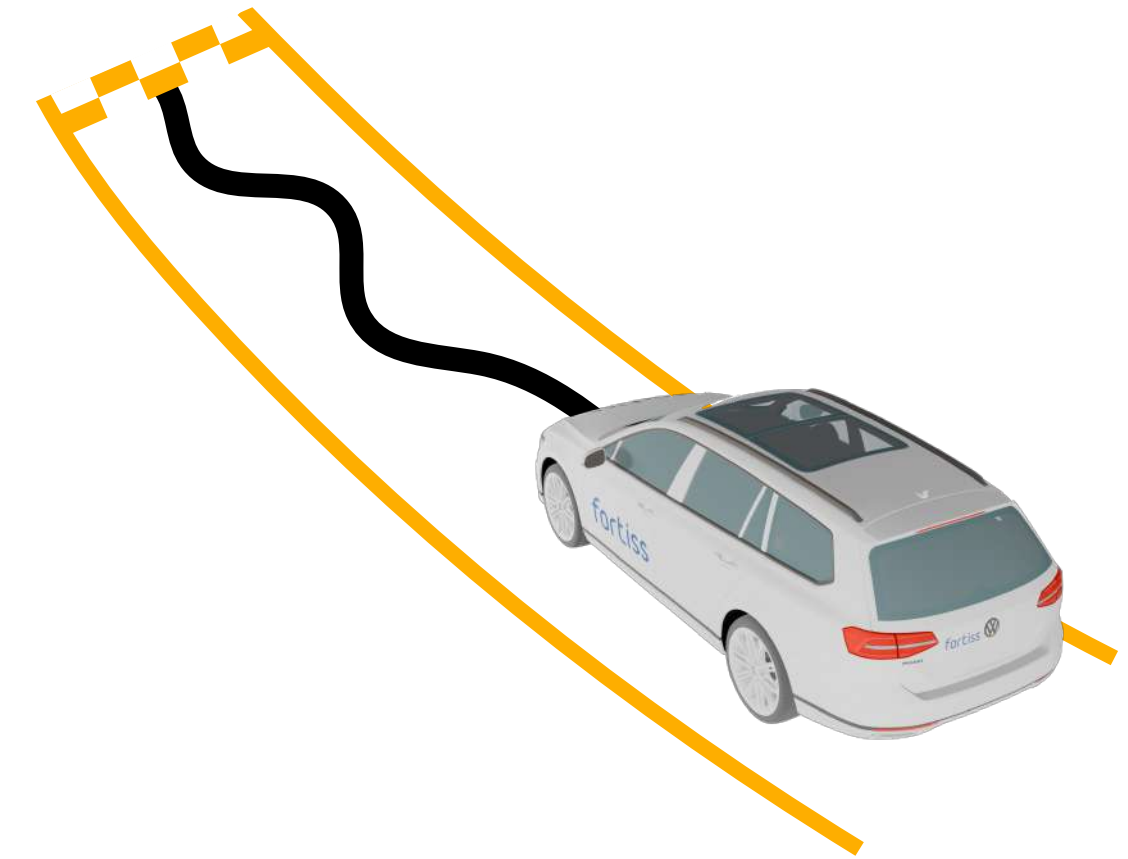
RQ1



Out of road

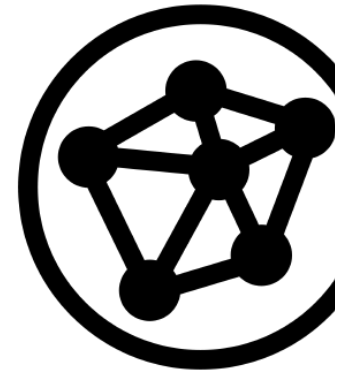


Out of time



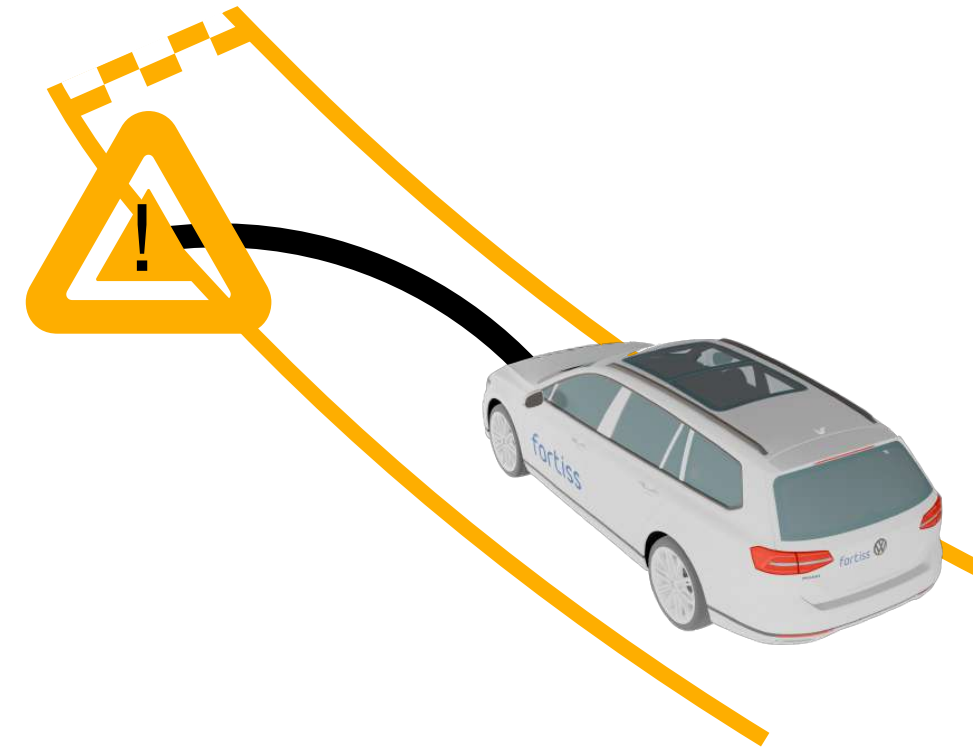
Jitter

RQ1

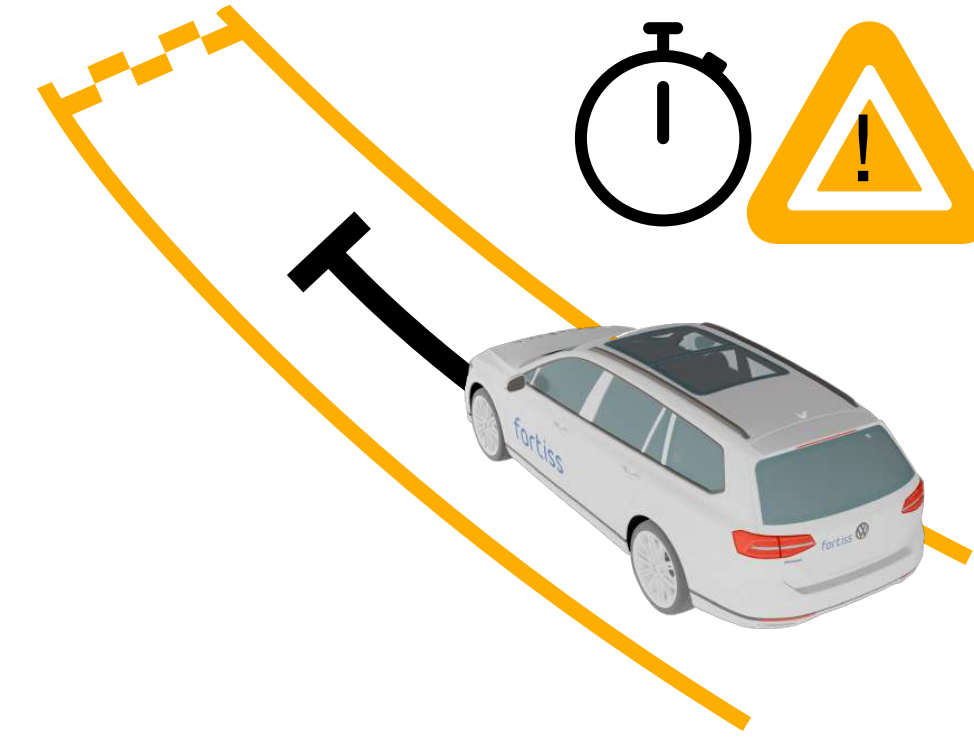


LK
ACC

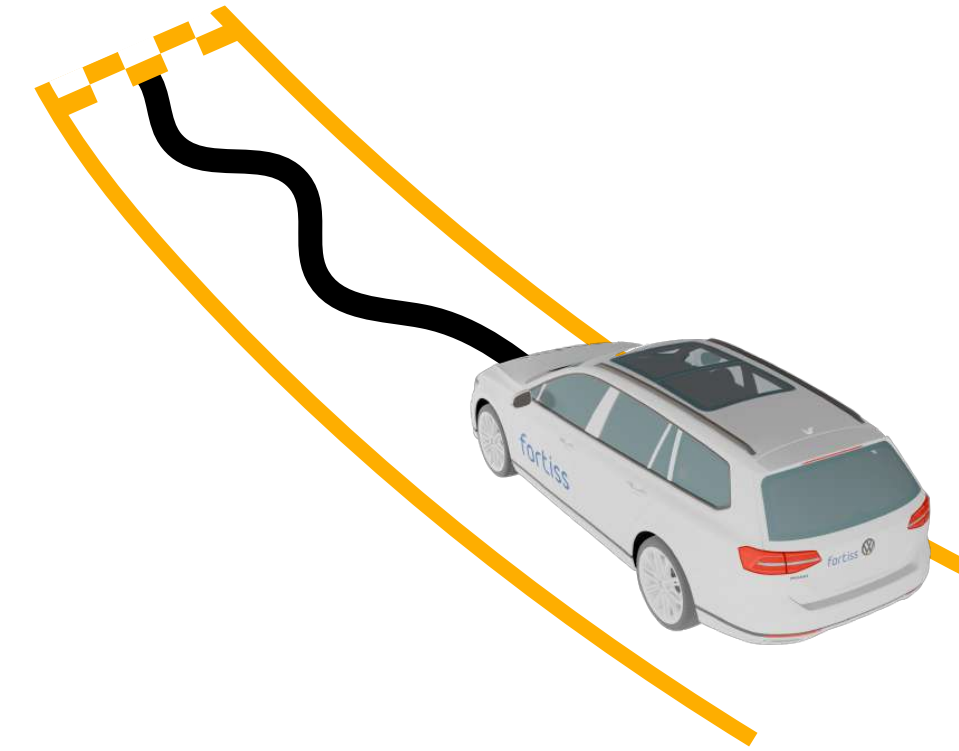
SIM2
Udacity



Out of road



Out of time



Jitter

False
color

↓ **76%**
Success
rate

+38

+0

Jpeg
artefacts

↓ **34%**
Success
rate

+1

+16

Saturation
decrease

↓ **48%**
Success
rate

+15

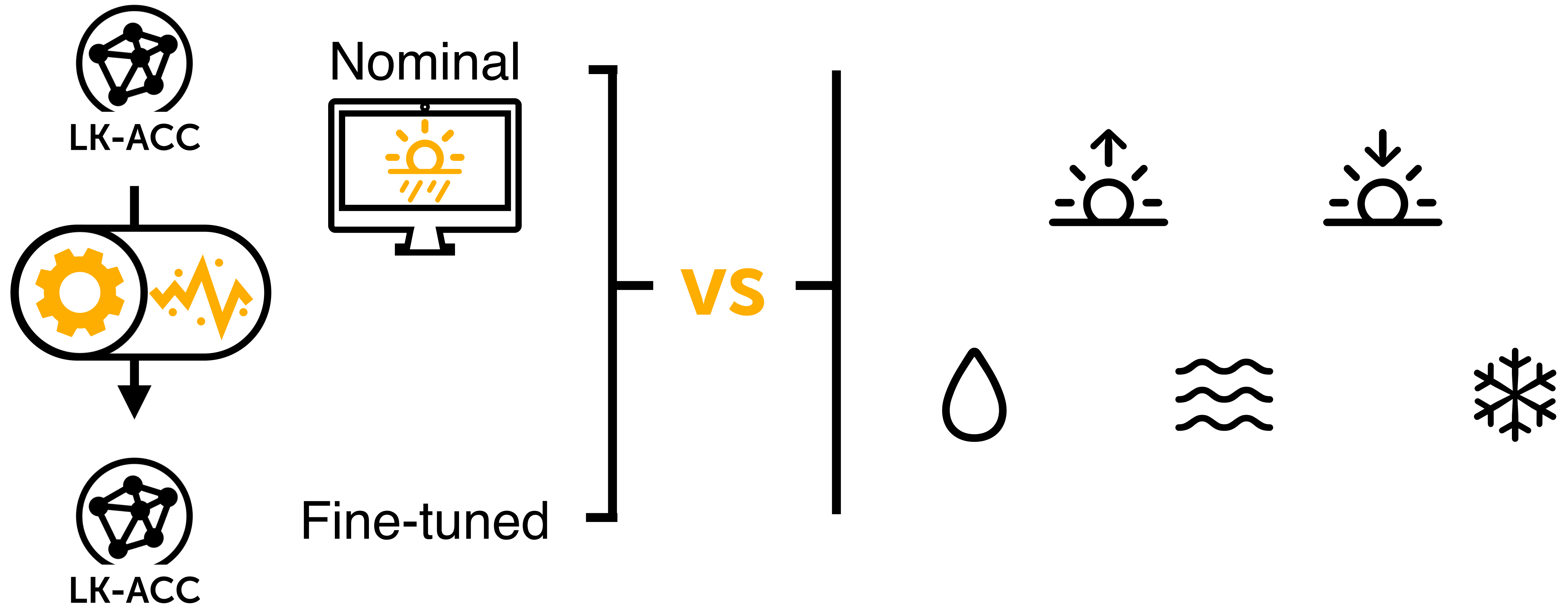
+9

↑ 530%

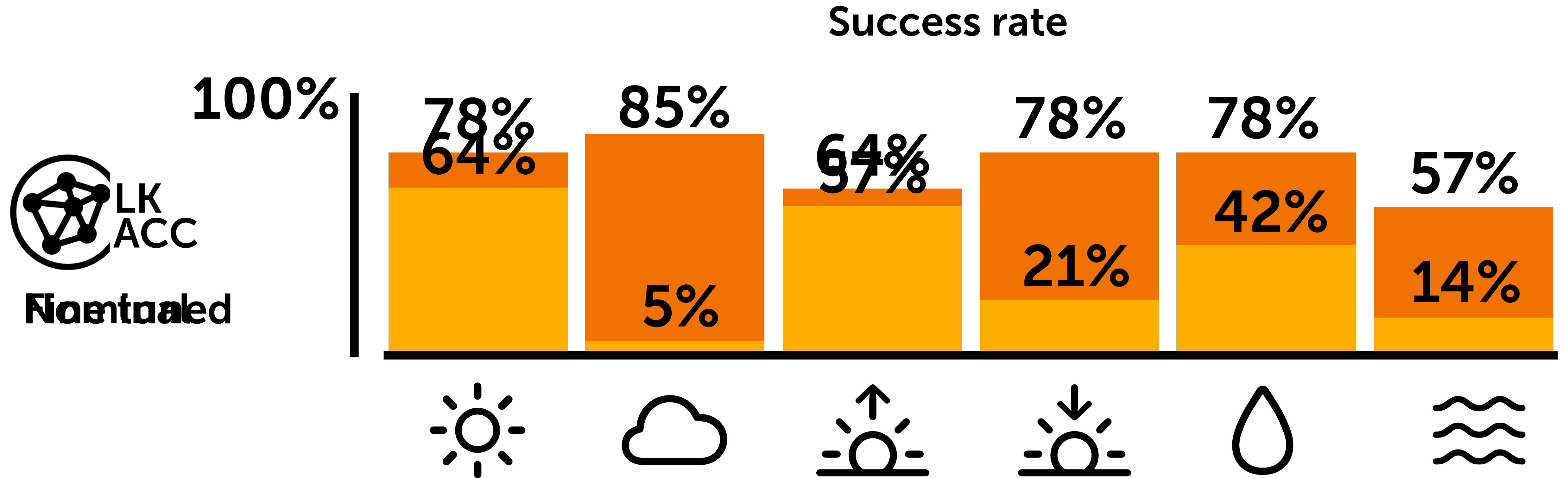
↑ 130%

↑ 621%

RQ2 Generalisation

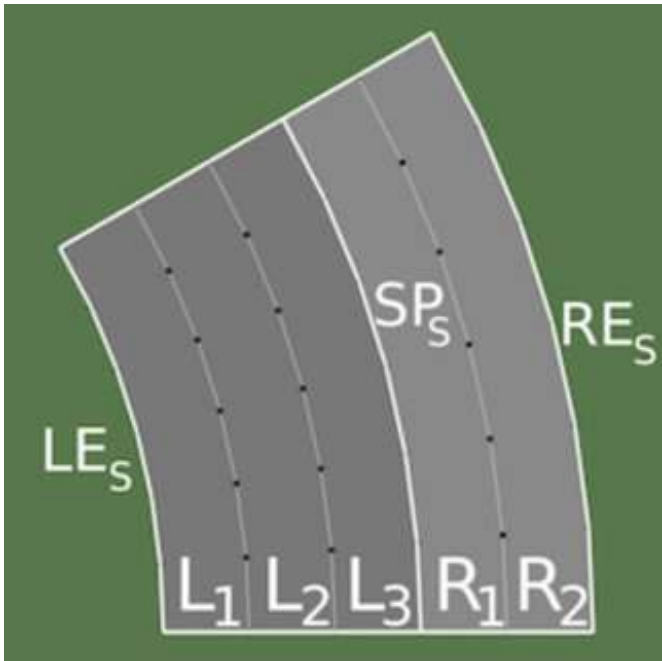


RQ2

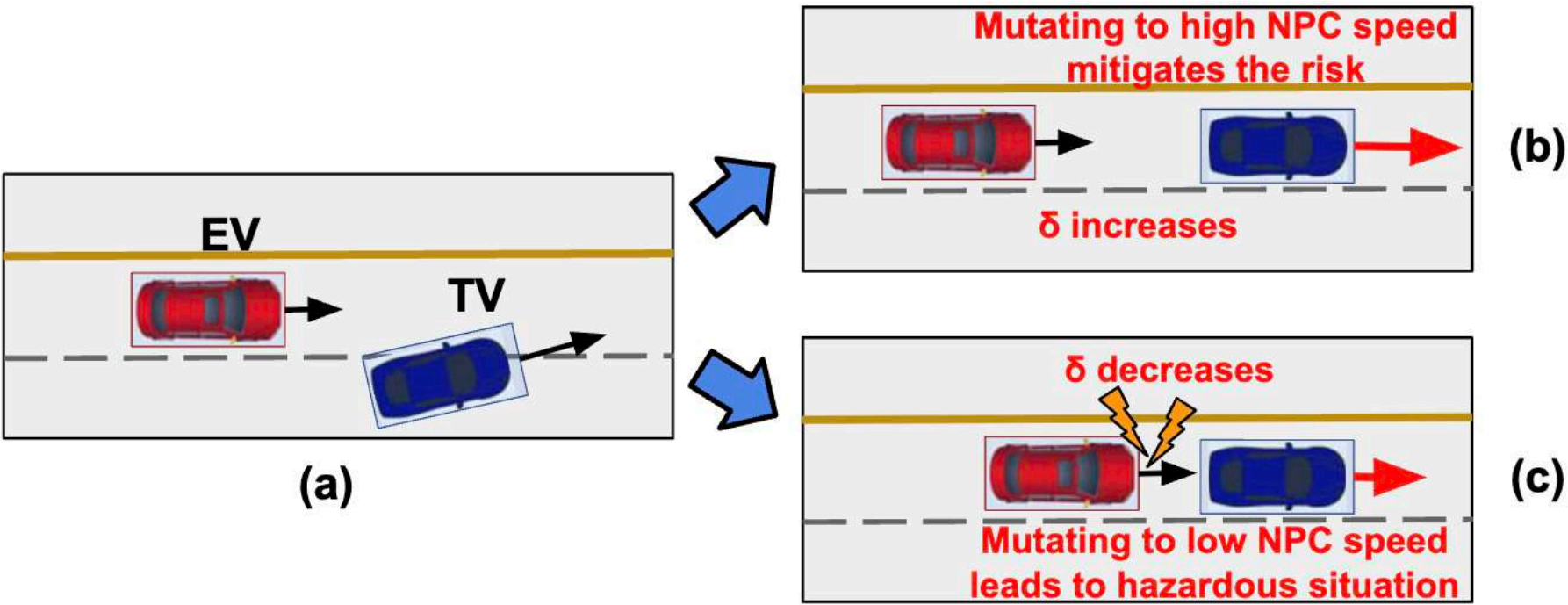


Testing for Generalization

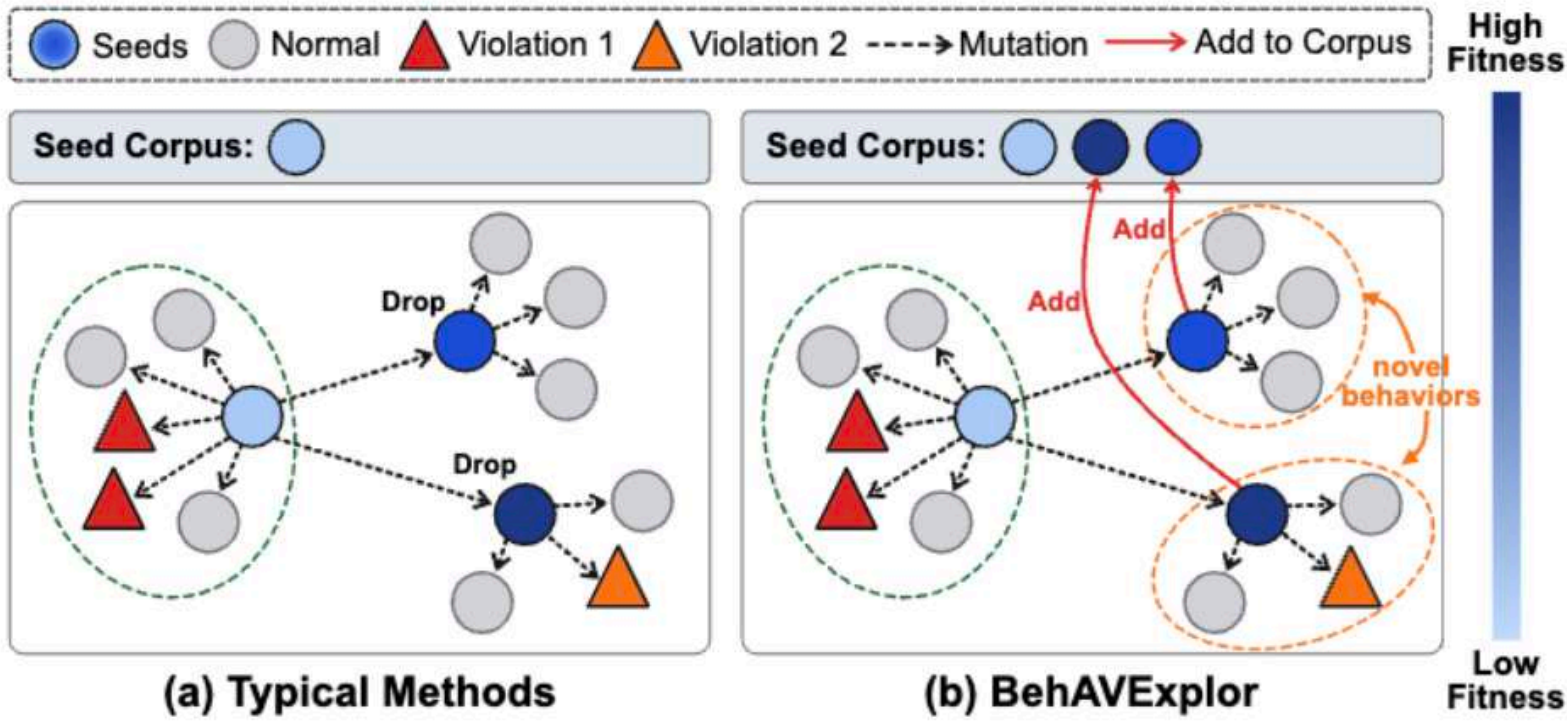
AsFault
Gambi et al., 2019



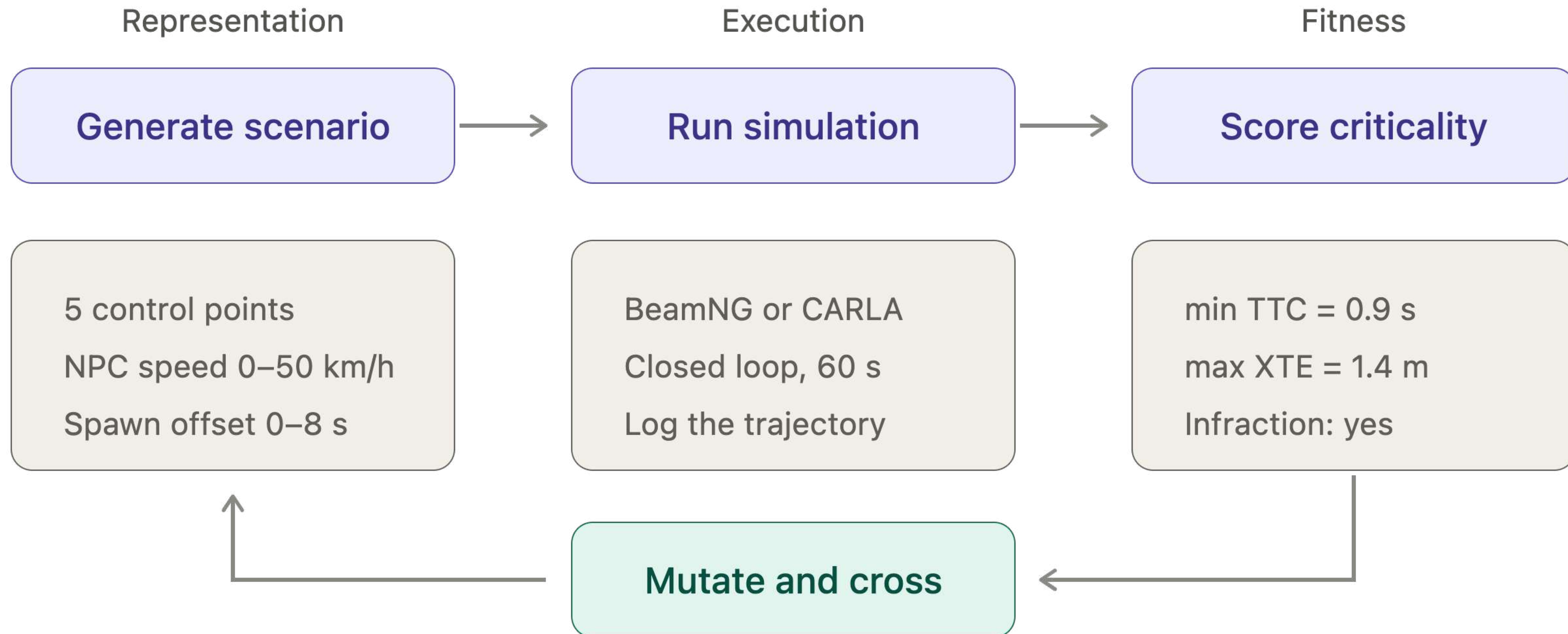
AVFuzzer
Li et al., 2020



BehAVEExplor
Cheng et al., 2023



Search-based Test Generation



Repeat until the budget is spent, then return the most critical scenarios found

OpenSBT

Open Search-based Testing



ADS & ADAS

- End-to end DNNs (level 2)
- Full AD stacks



Critical Test Scenarios Identification



High fidelity sim-based Validation

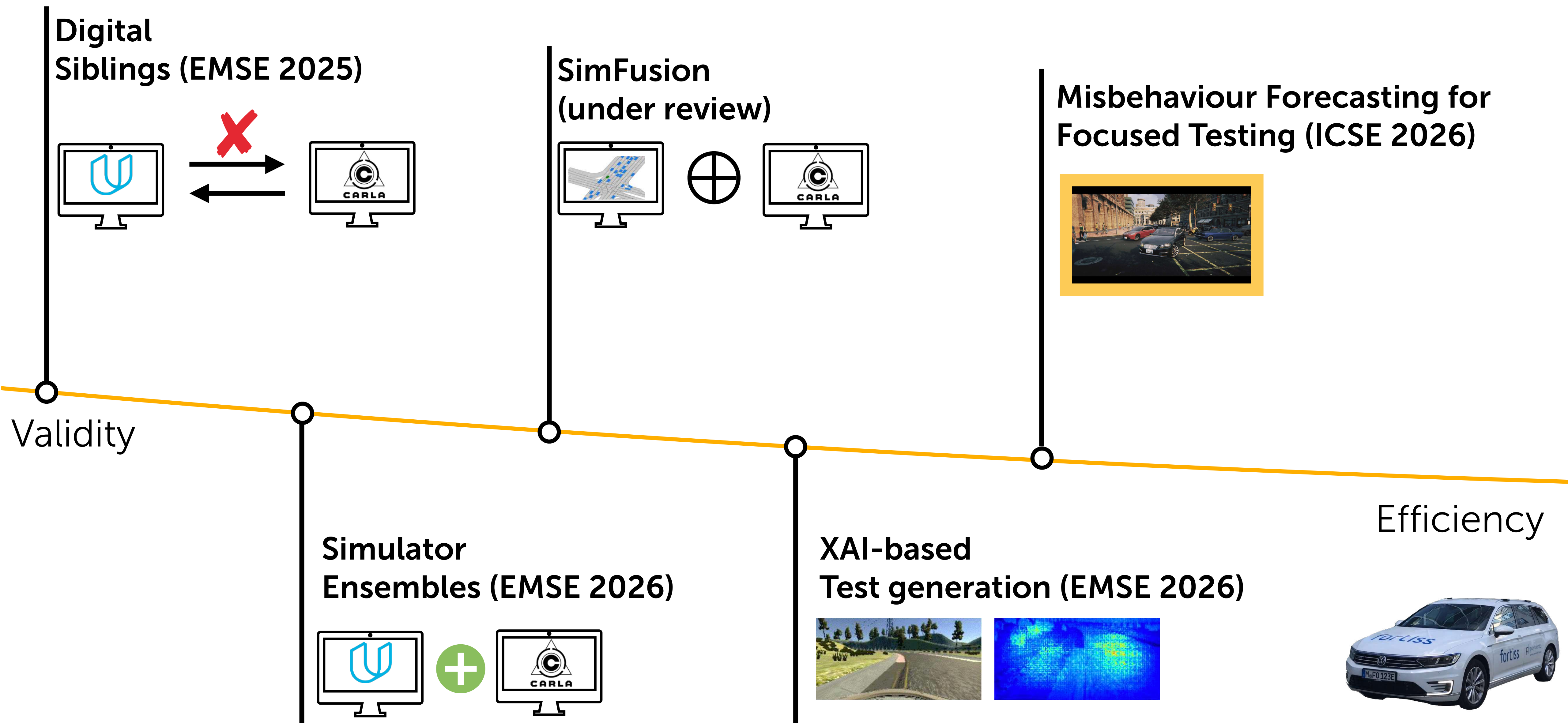


<https://git.fortiss.org/opensbt>



Sorokin et al. OpenSBT: A Modular Framework for Search-based Testing of Automated Driving Systems.

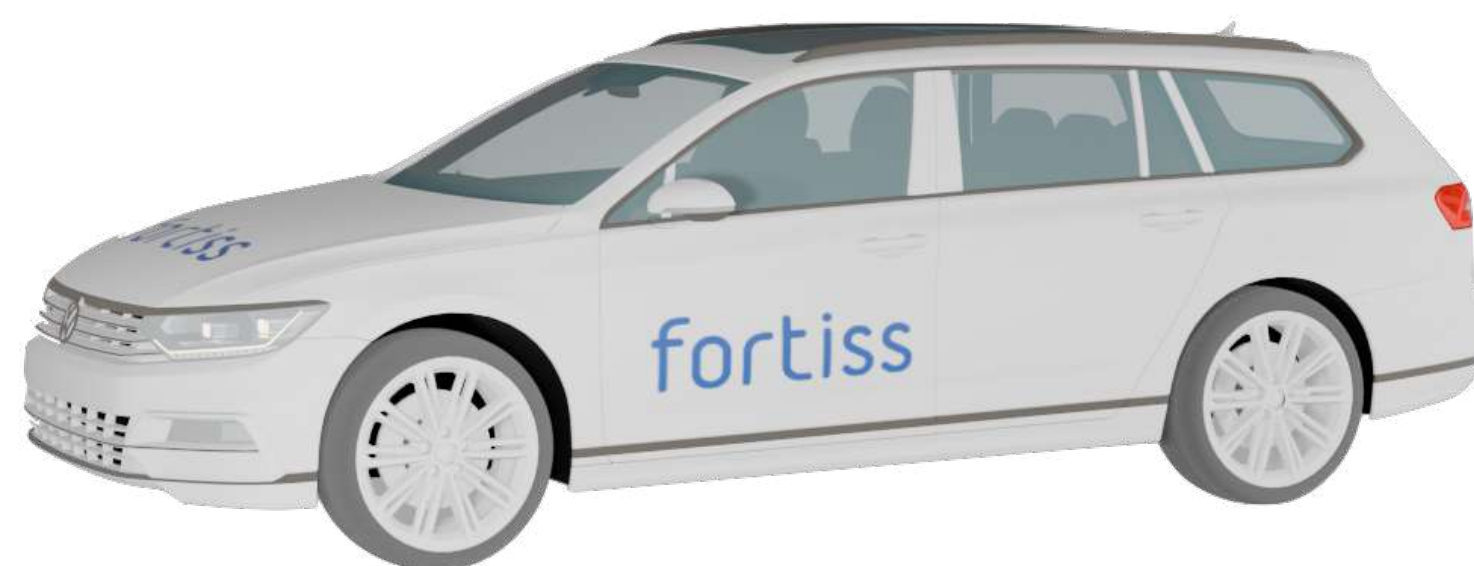
Automated Test Generation



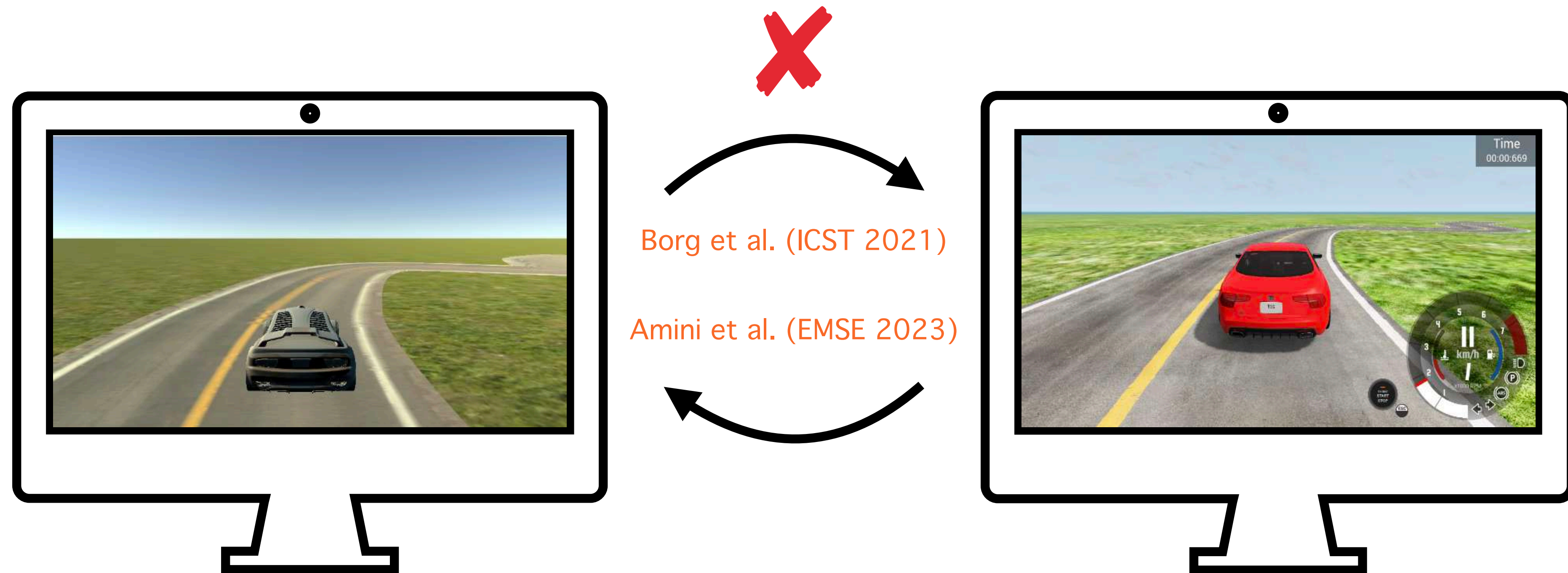
How to increase validity?

Real vehicles?

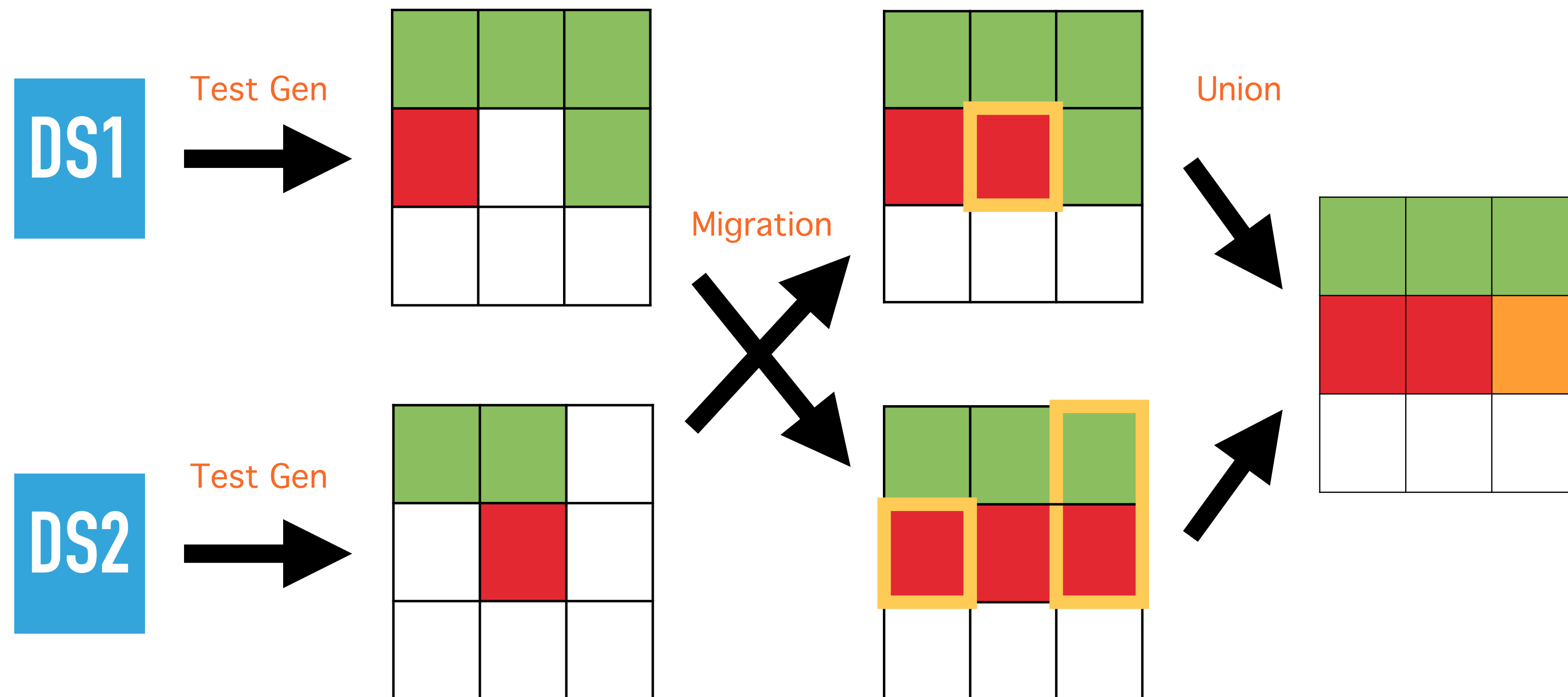
Digital Twins



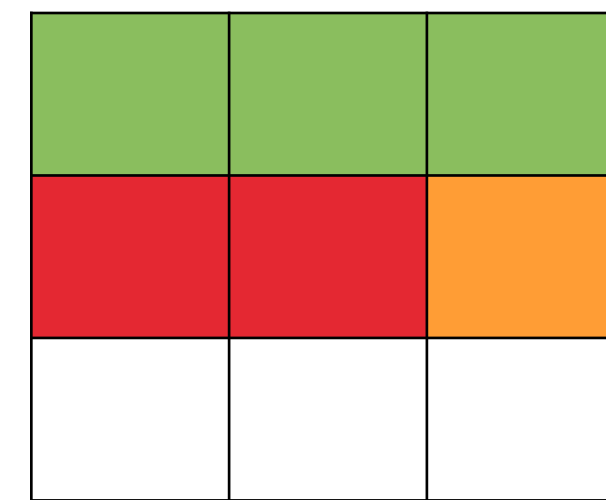
General-purpose simulators



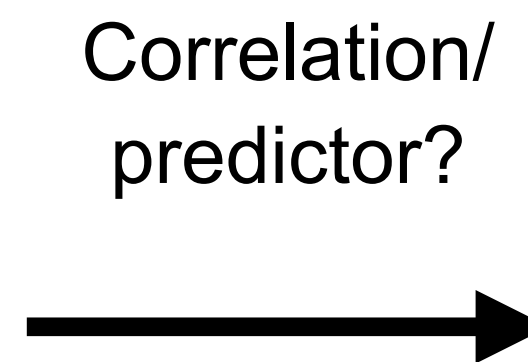
Digital Siblings



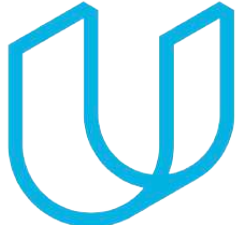
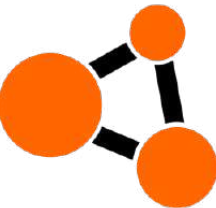
Digital siblings vs digital twin

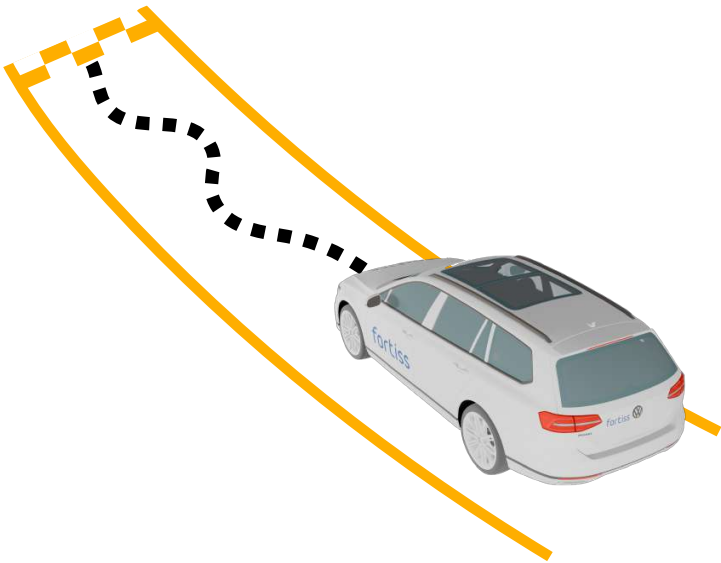


Digital siblings

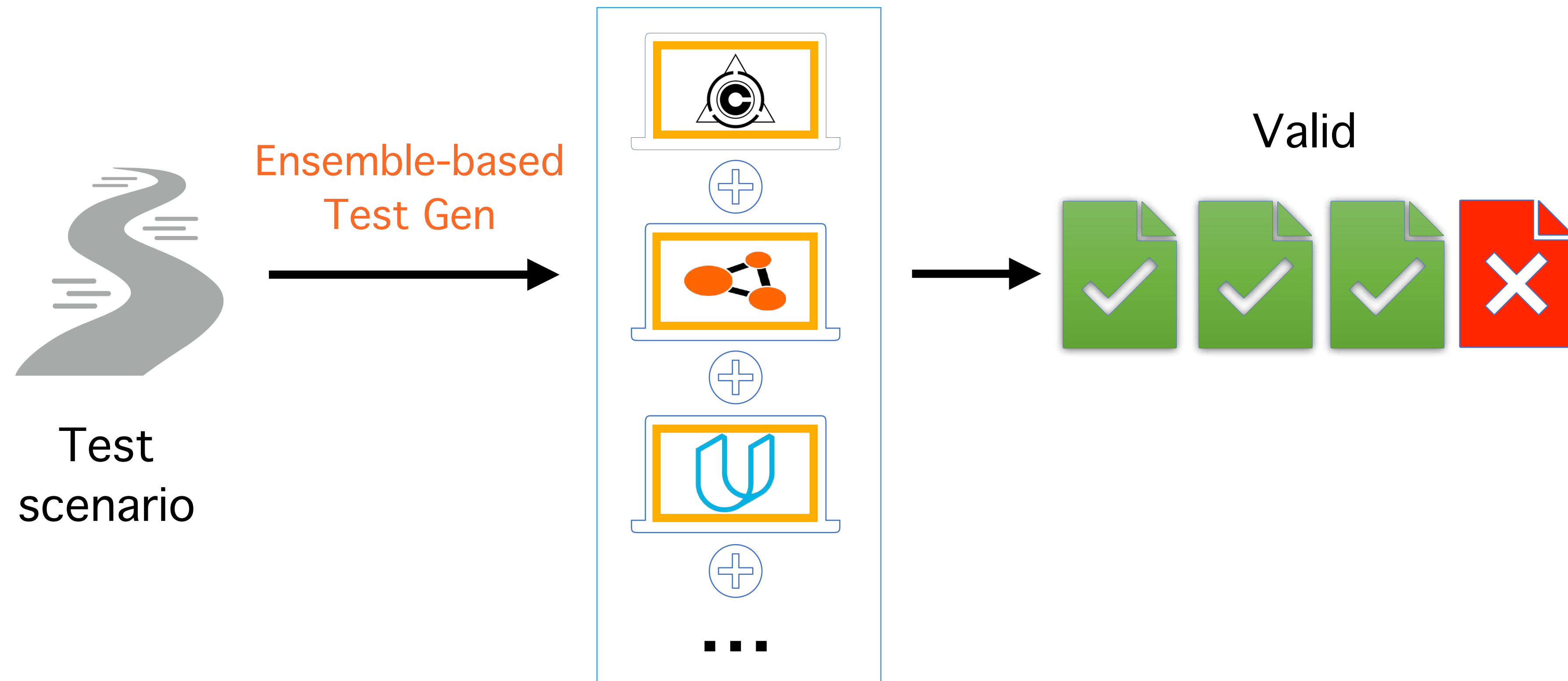


Digital twin

	Dave-2	Chauffeur	Epoch
 vs DT	0.459	0.740	0.392
 vs DT	0.496	0.791	0.541
   vs DT	0.659	0.940	0.594

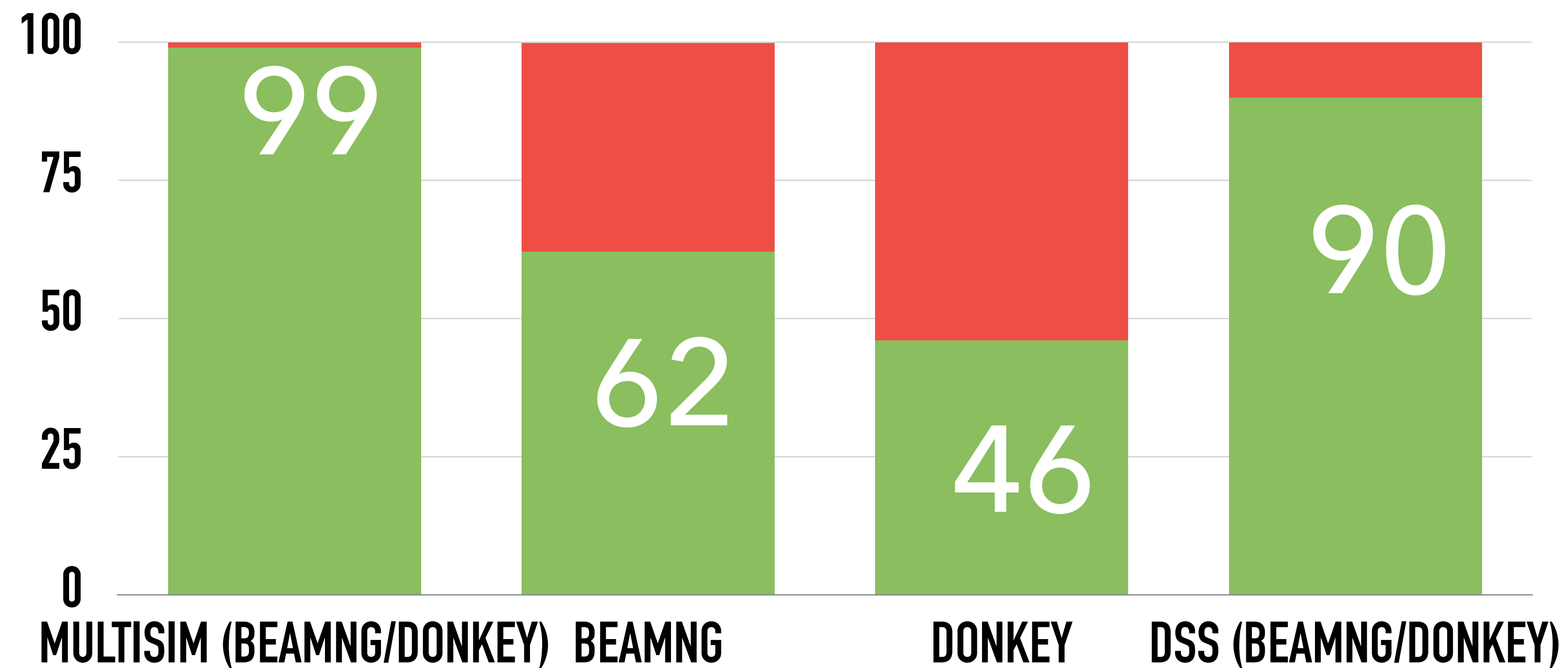


Simulators ensemble

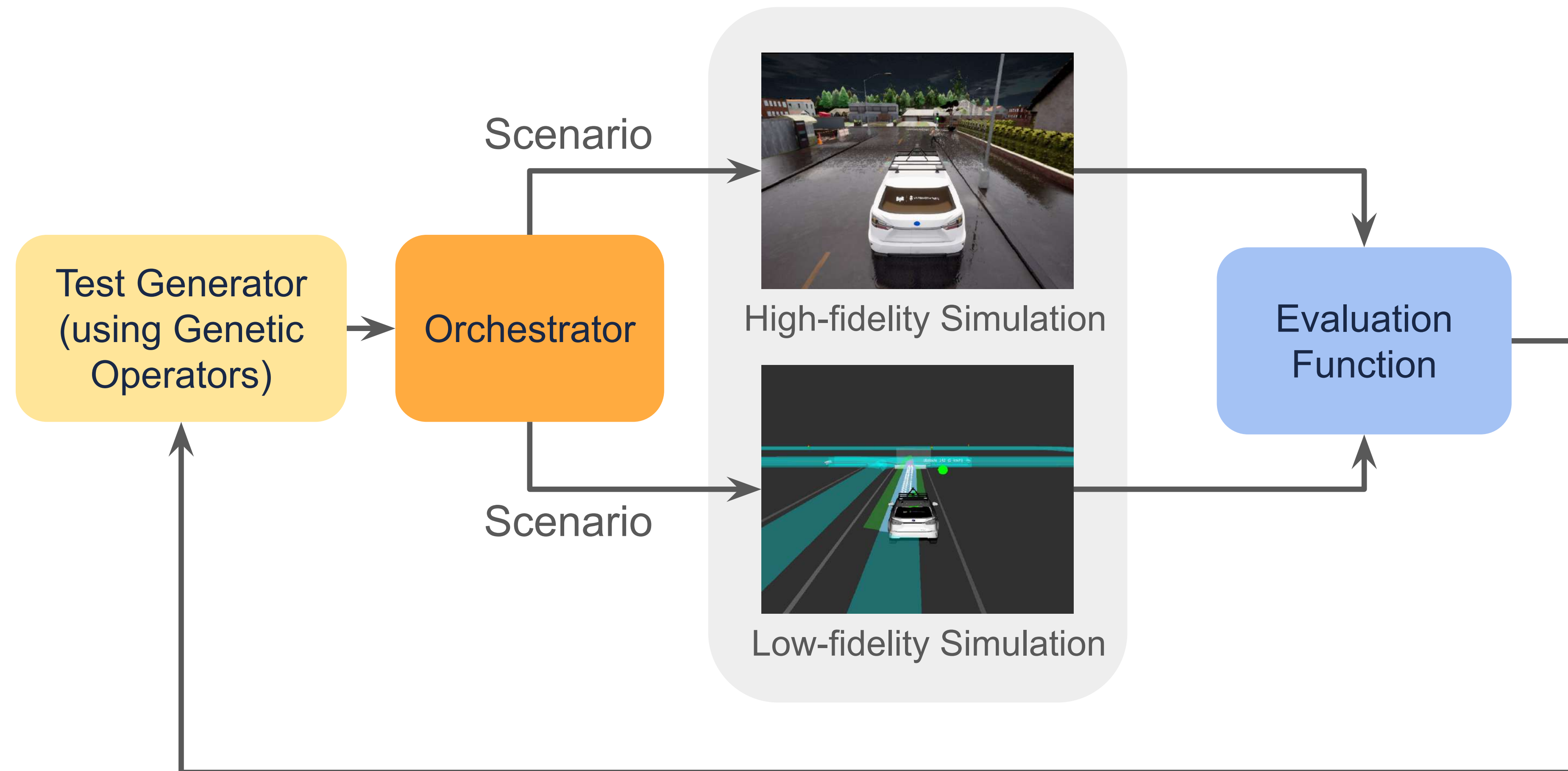


L. Sorokin, M. Biagiola, A. Stocco. Simulator Ensembles for Trustworthy Autonomous Driving Systems Testing. Empirical Software Engineering, 2026

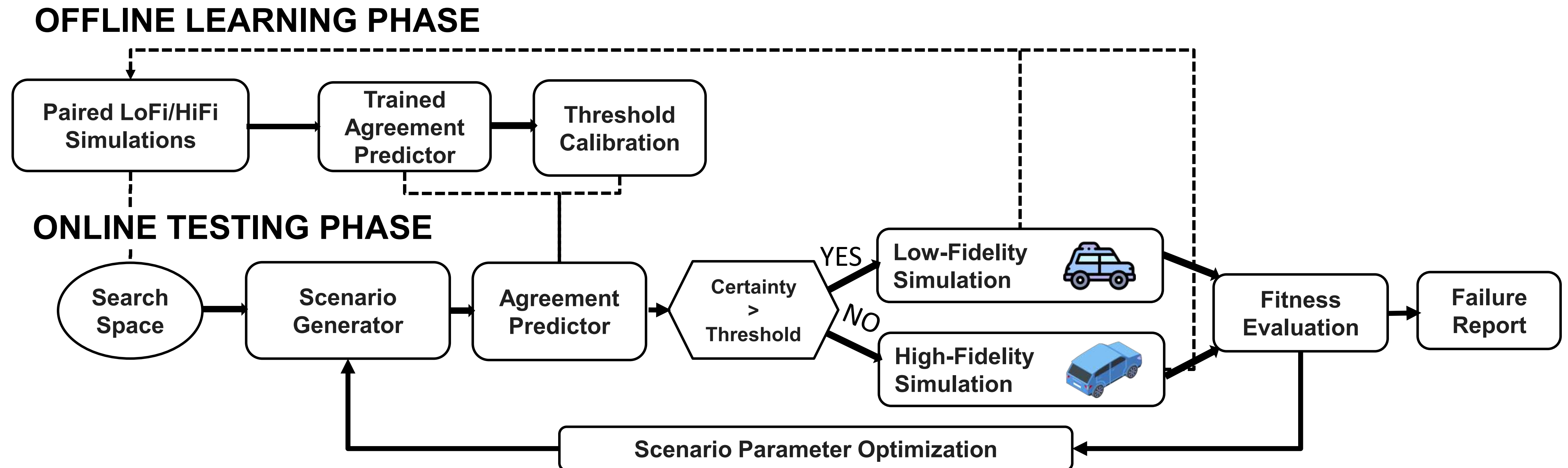
Increased failure validity



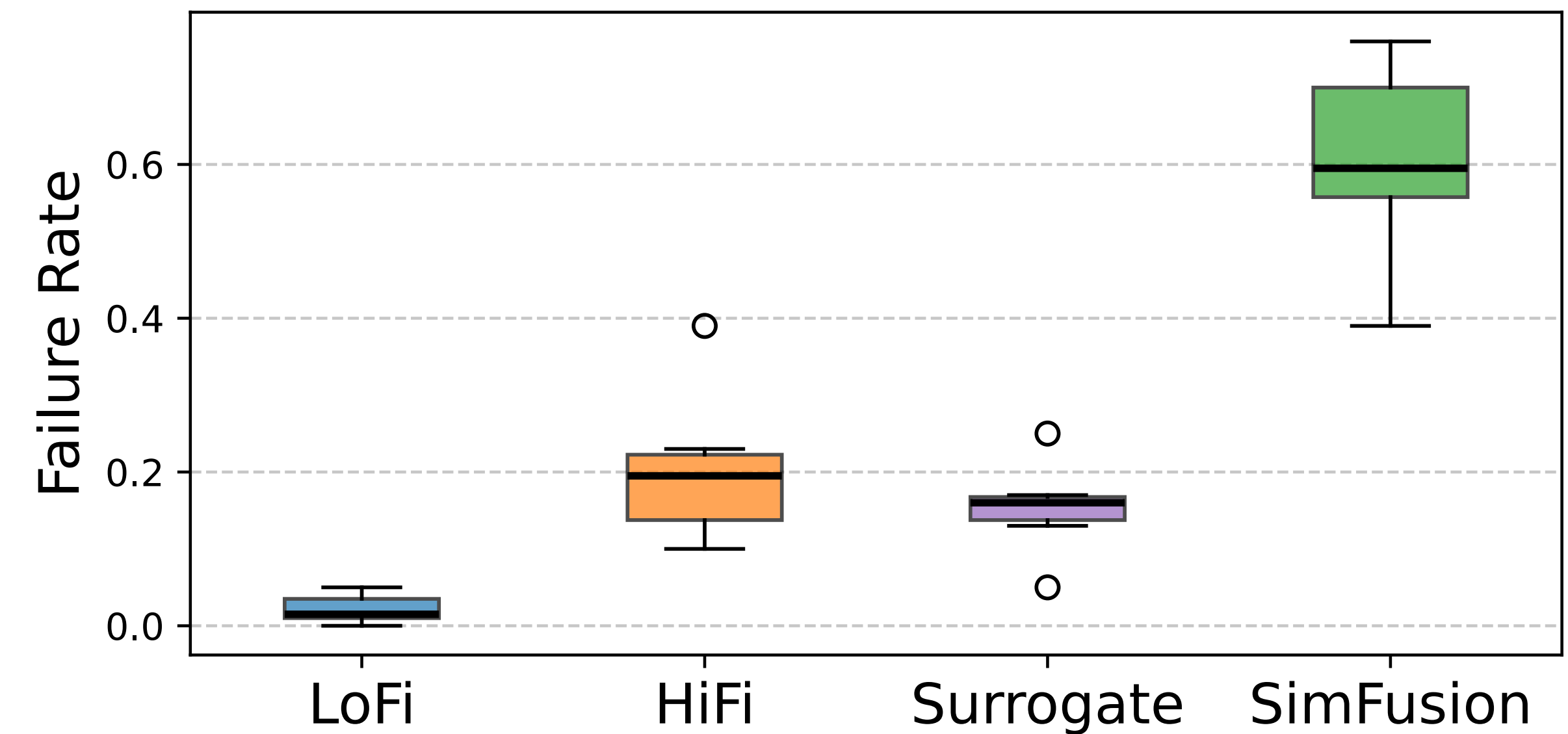
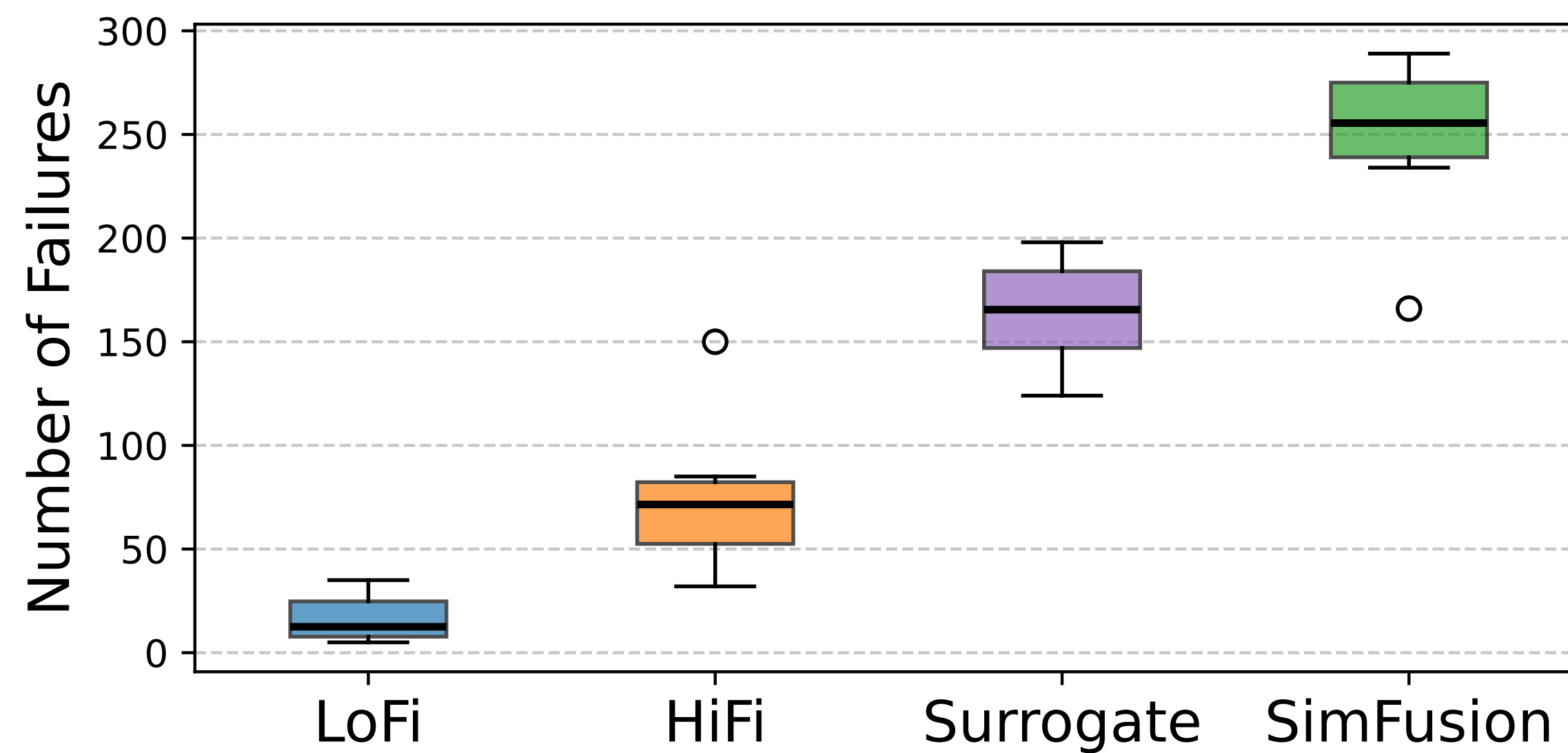
Multi-fidelity simulation



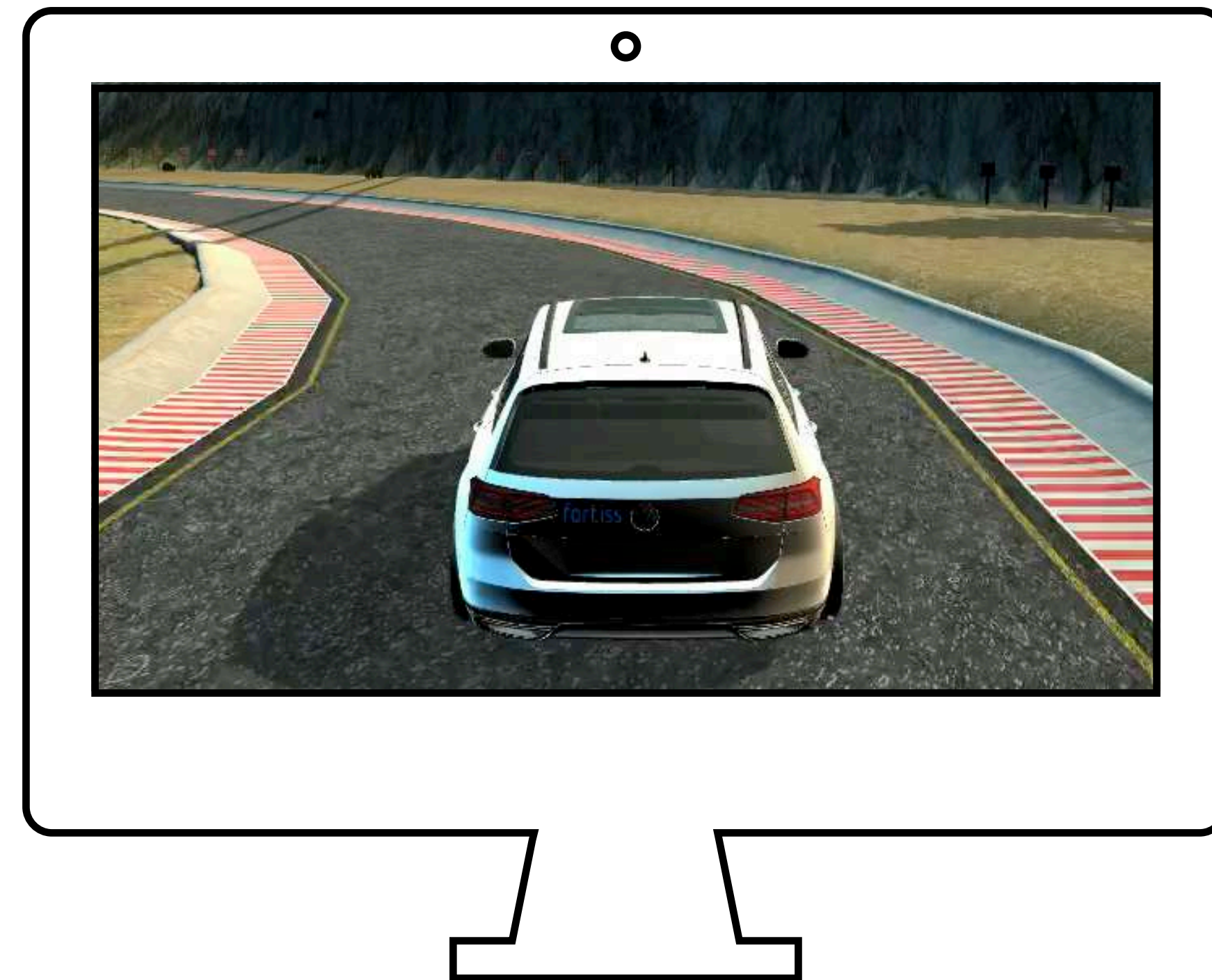
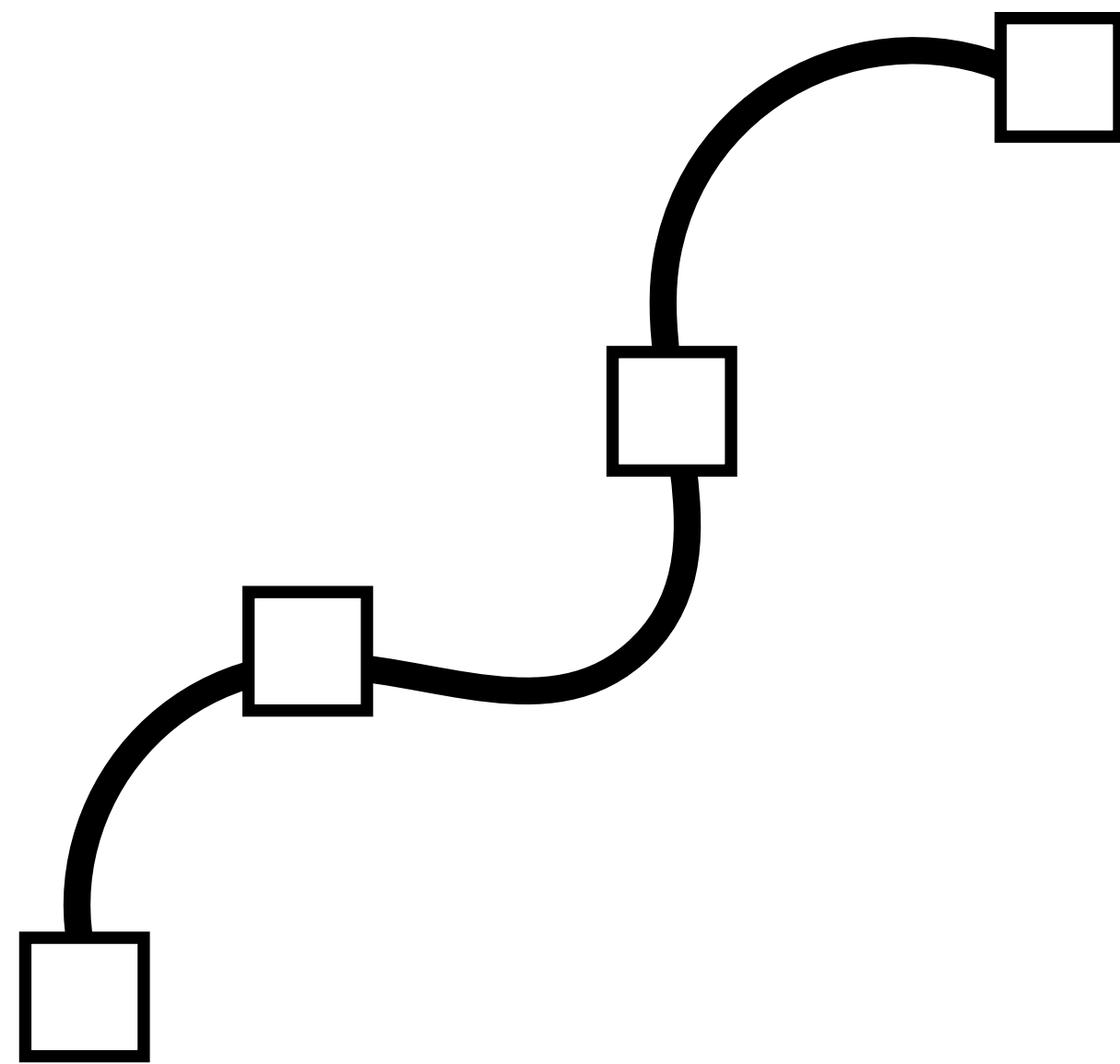
Multi-fidelity simulation



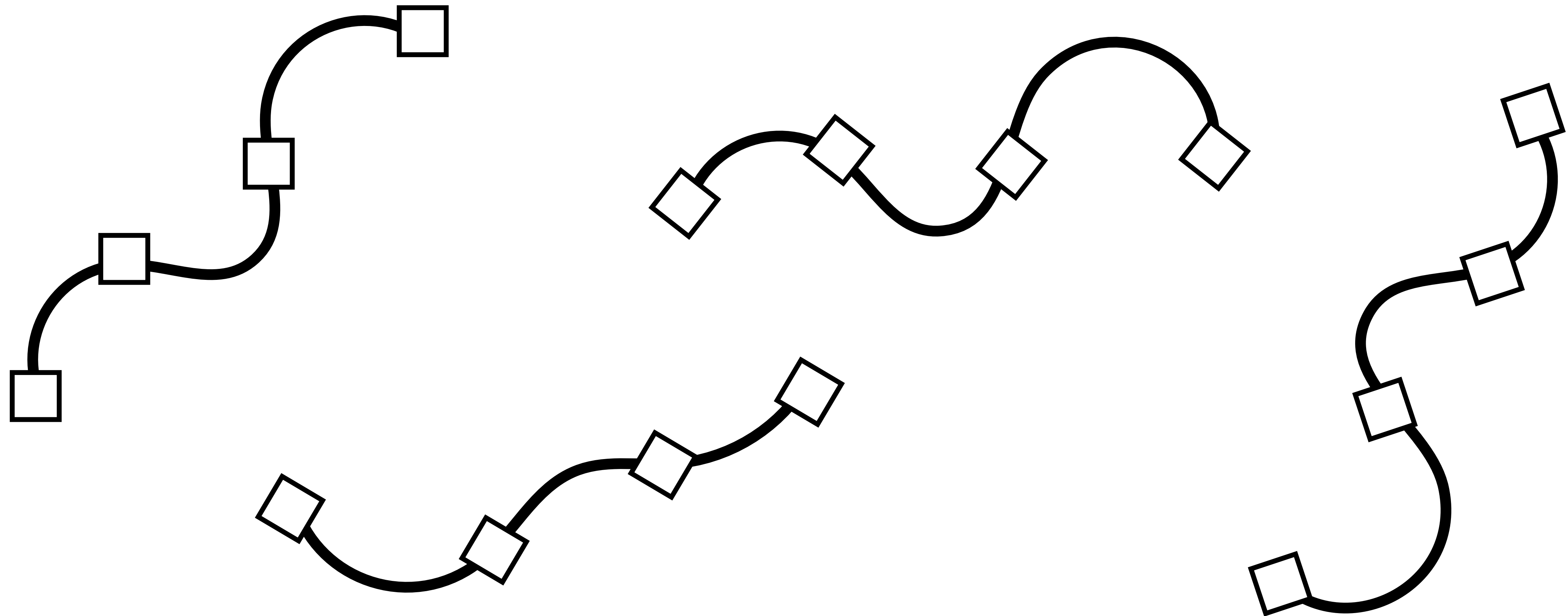
Multi-fidelity simulation



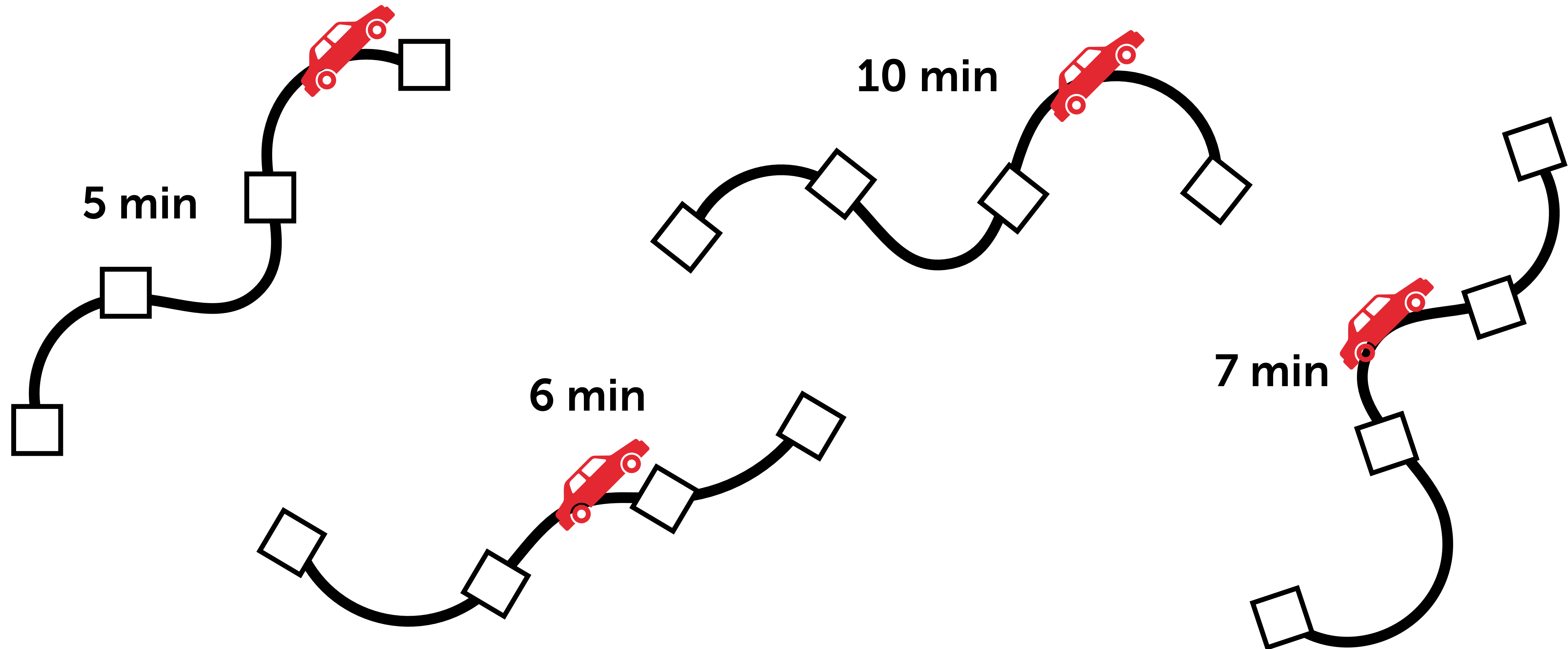
How to increase efficiency?



How to increase efficiency?

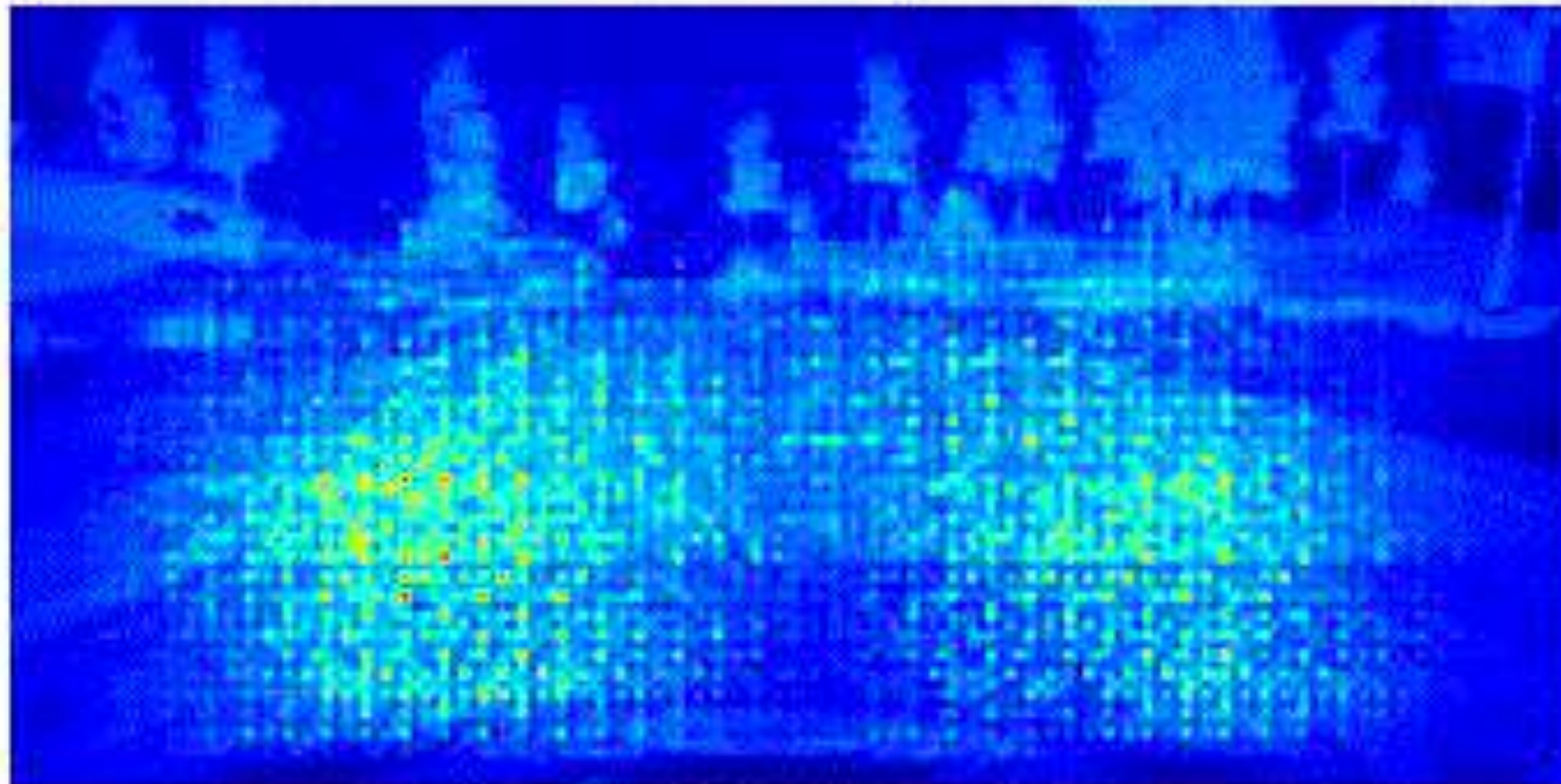


How to increase efficiency?



Nominal

XAI Image

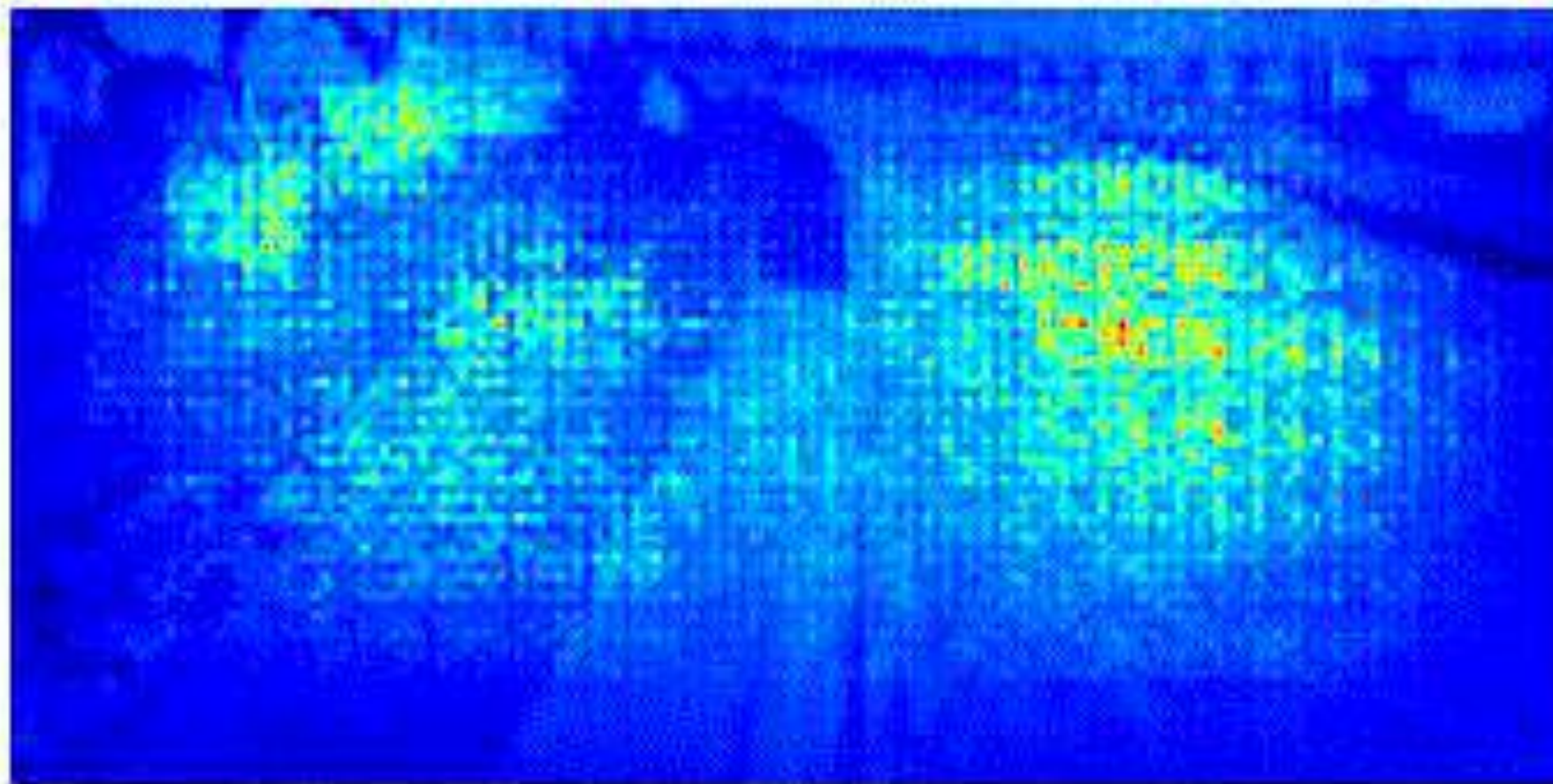


Original Image



Uncertain/Unexpected

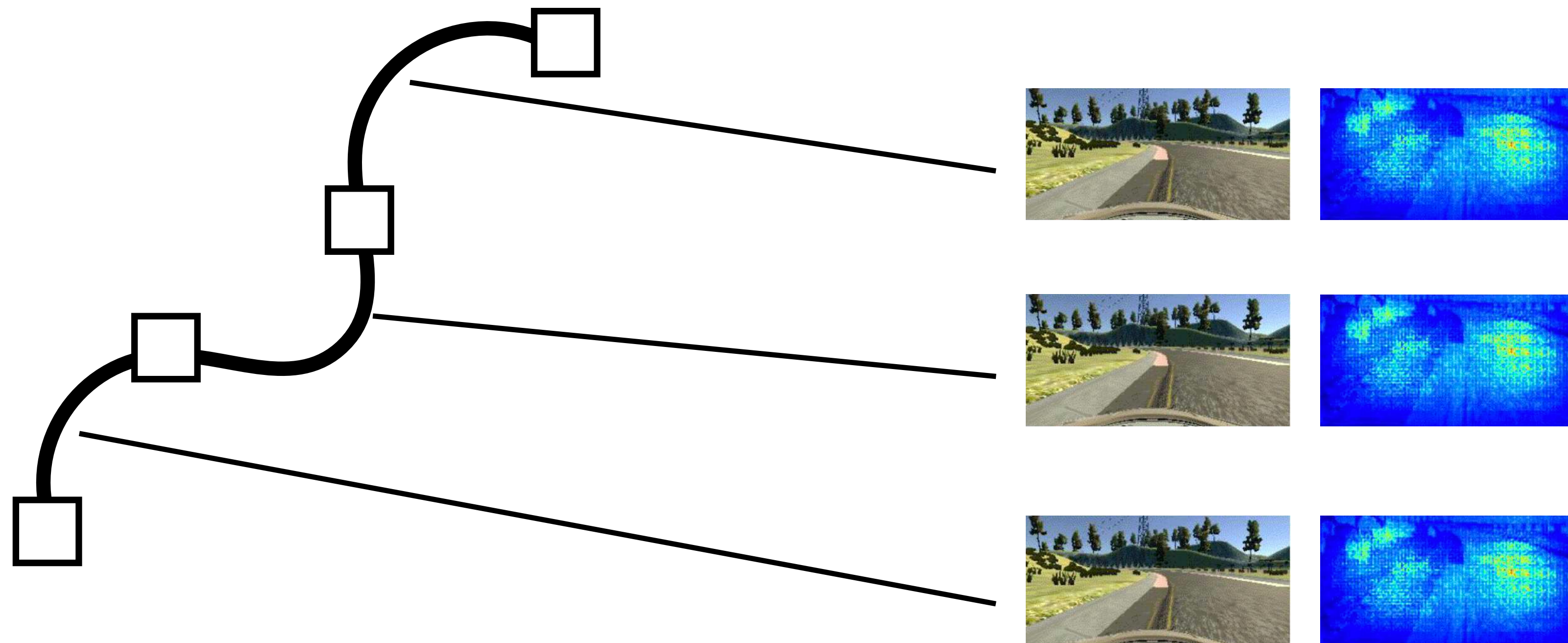
XAI Image



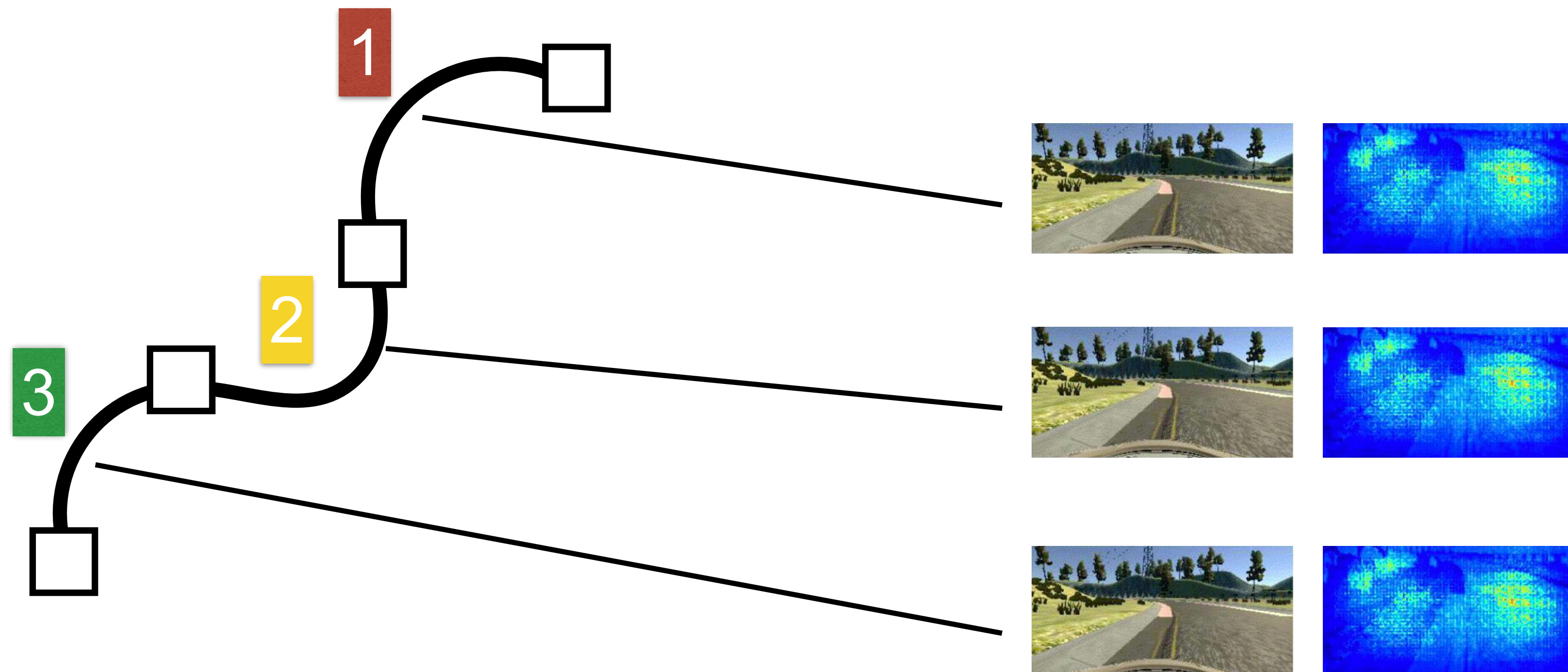
Original Image



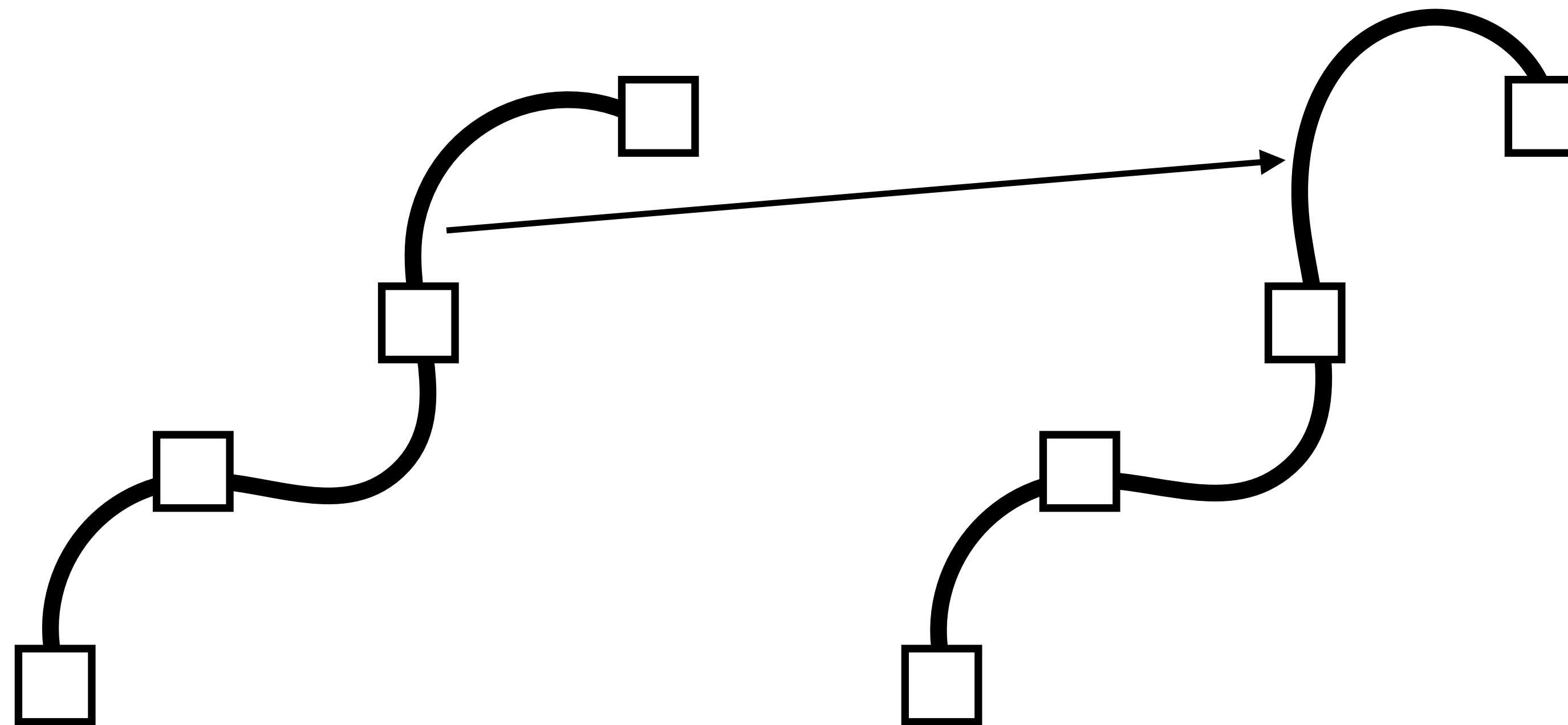
XAI-guided Test Generation



XAI-guided Test Generation

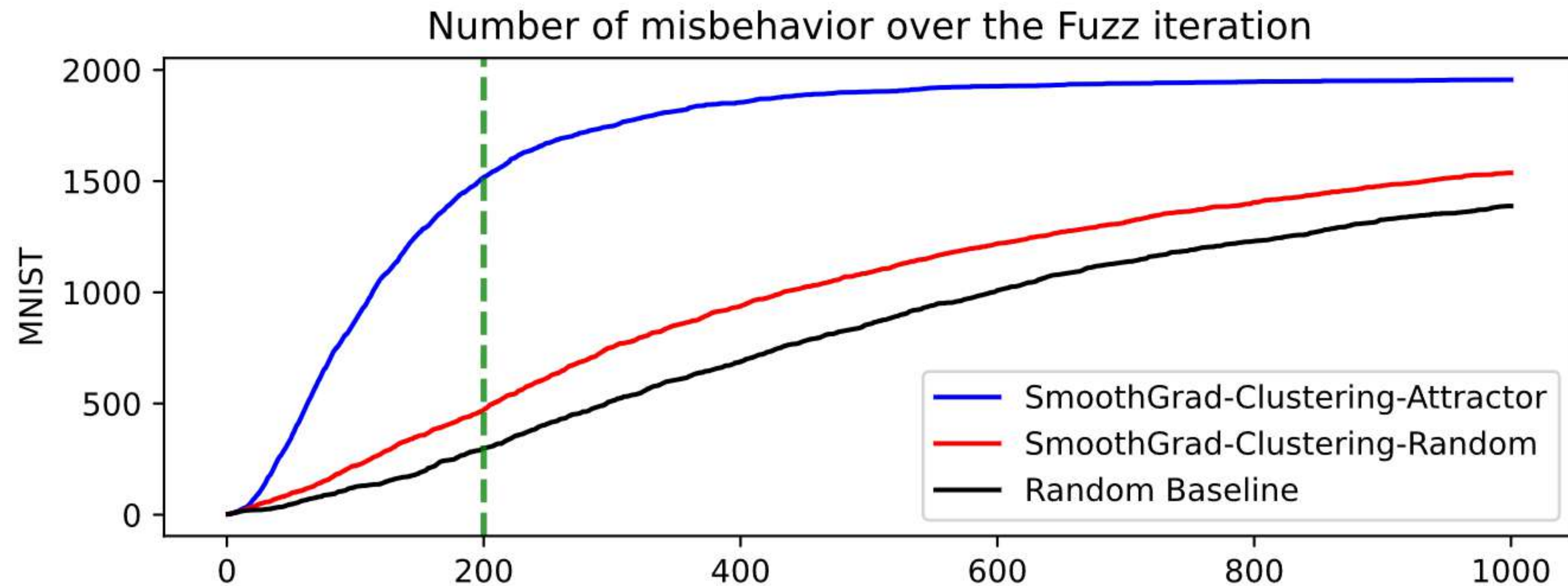


XAI-guided Test Generation



Result

Effectiveness (200% more),
Performance (up to x7)



X. Chen, M. Biagiola, M. D'Amorim, A. Stocco. XMutant: XAI-based Fuzzing for Deep Learning Systems. Empirical Software Engineering, 2025

Focused Test Generation for ADS



Pass Test



Pass Test

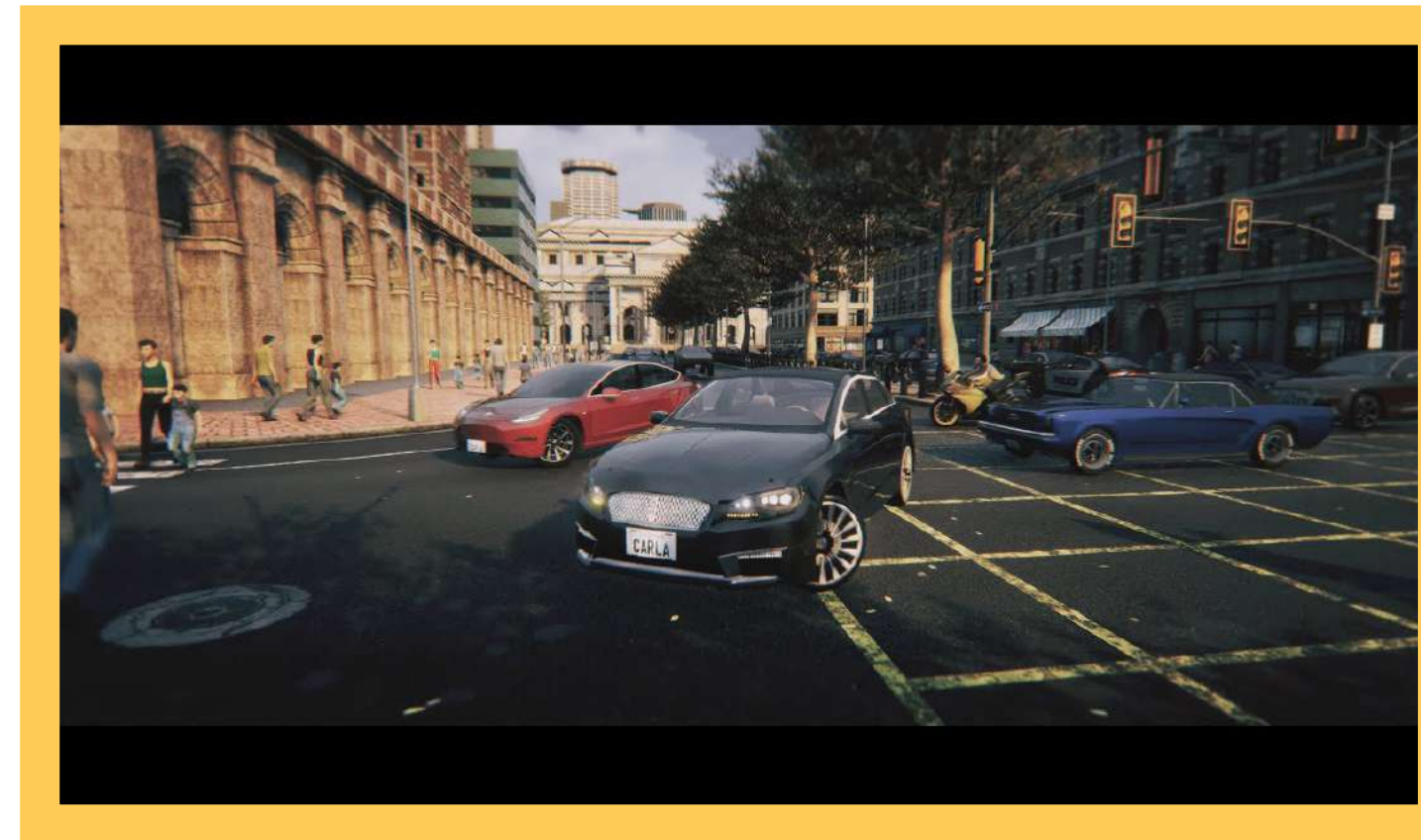


Fail Test

Focused Test Generation for ADS



Pass Test

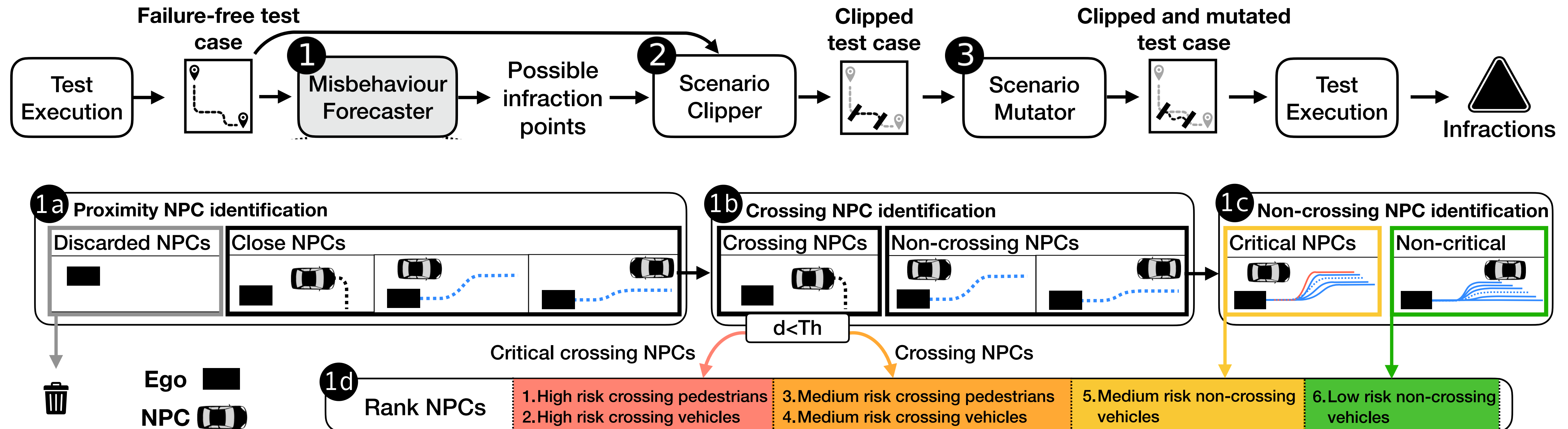


Near-miss (Fail) Test

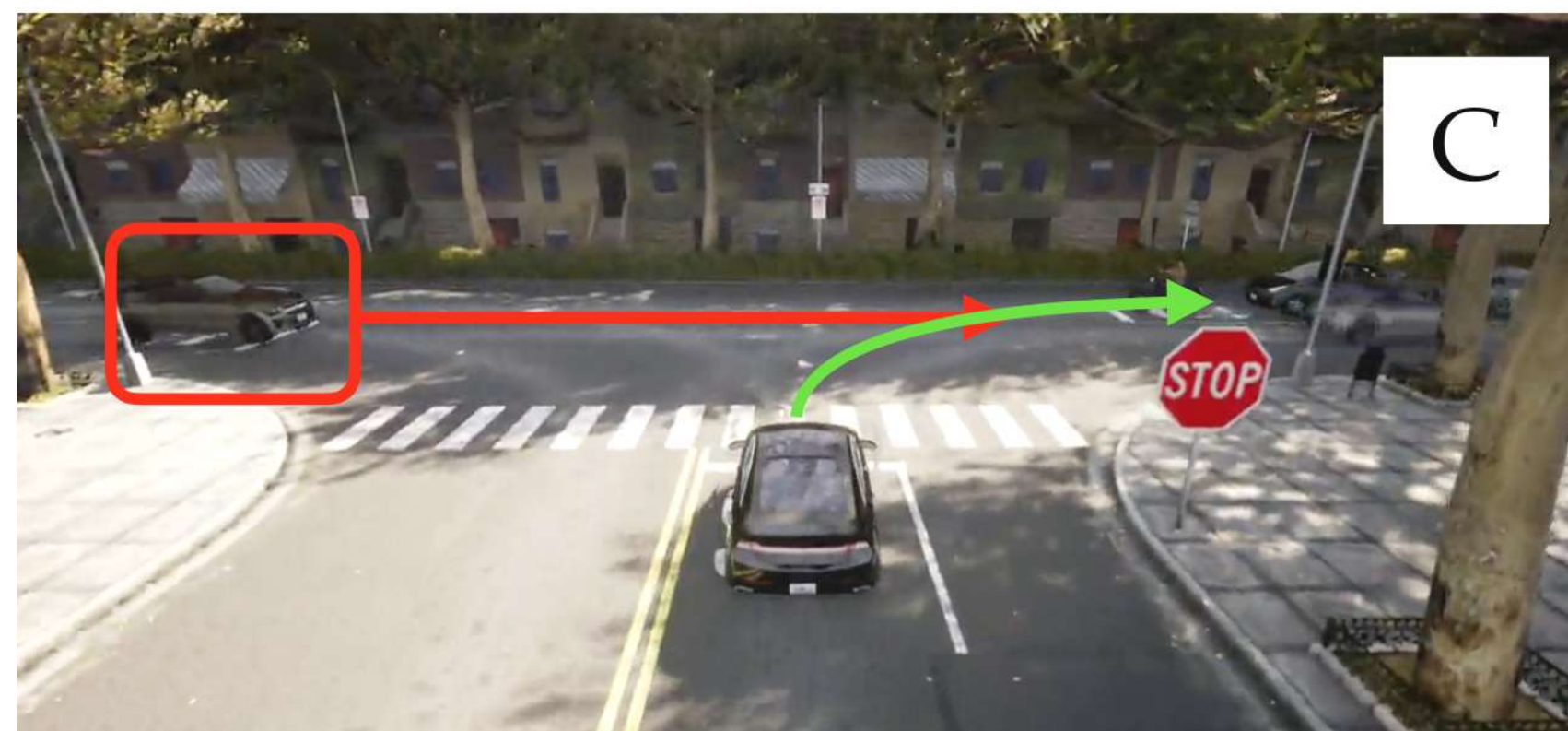
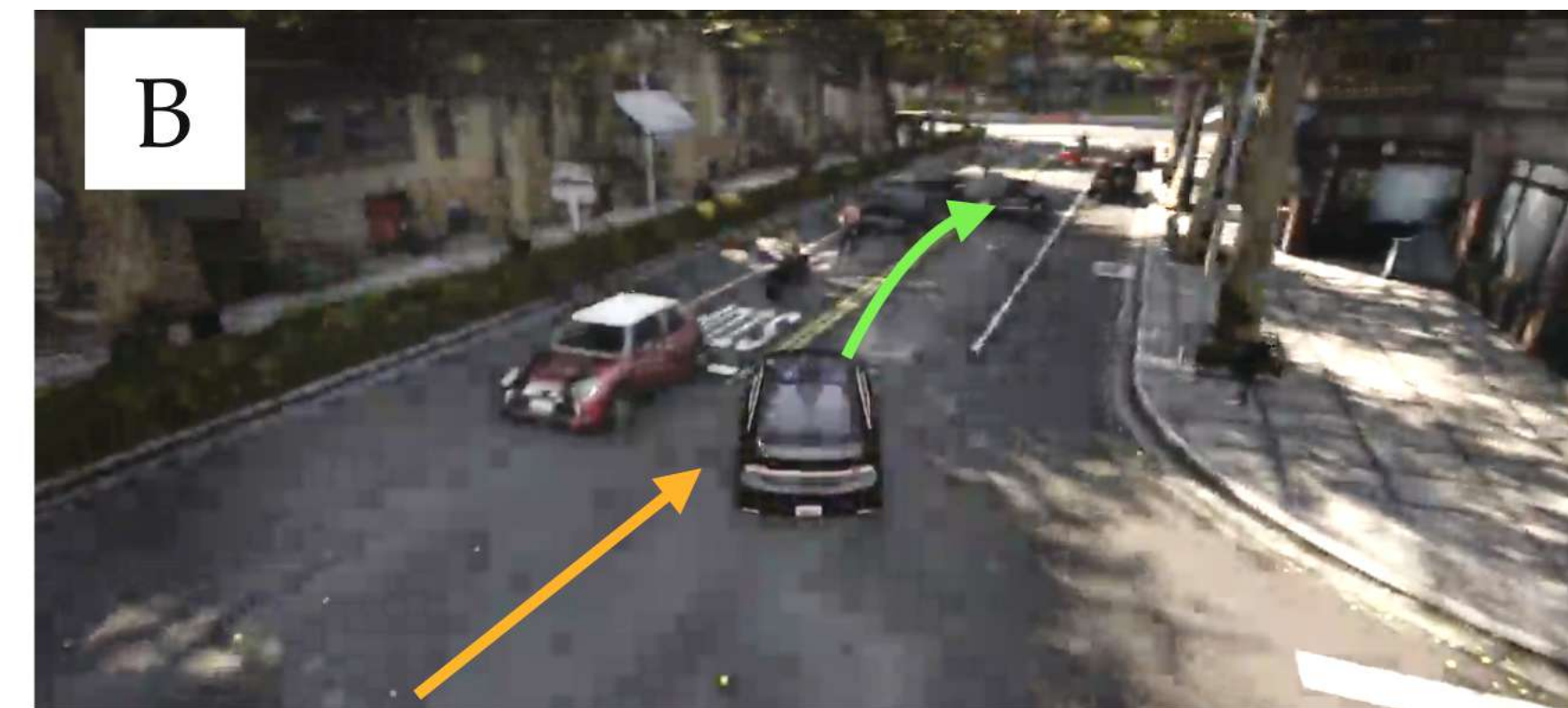
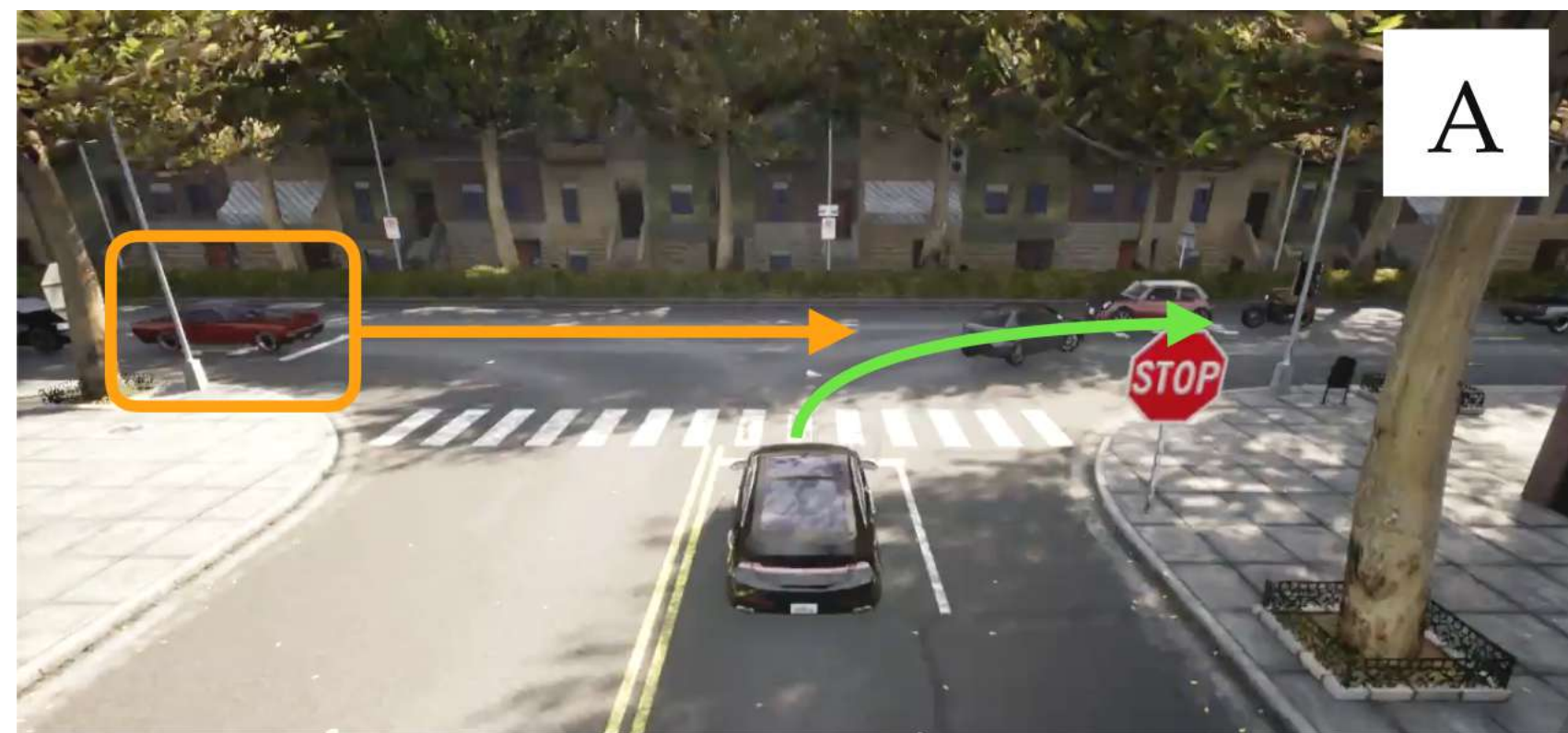


Fail Test

Focused Test Generation for ADS



Focused Test Generation for ADS

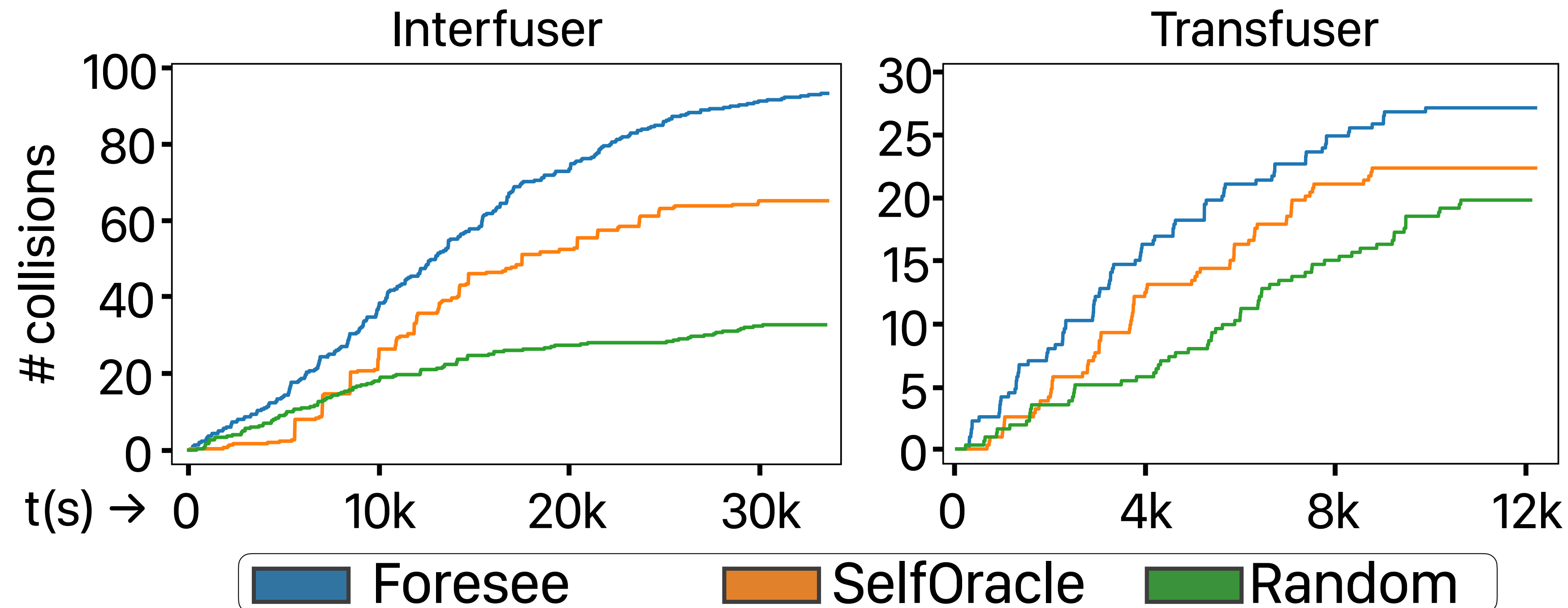


→ EGO vehicle

→ Nominal NPC

→ Fuzzed NPC

Focused Test Generation for ADS

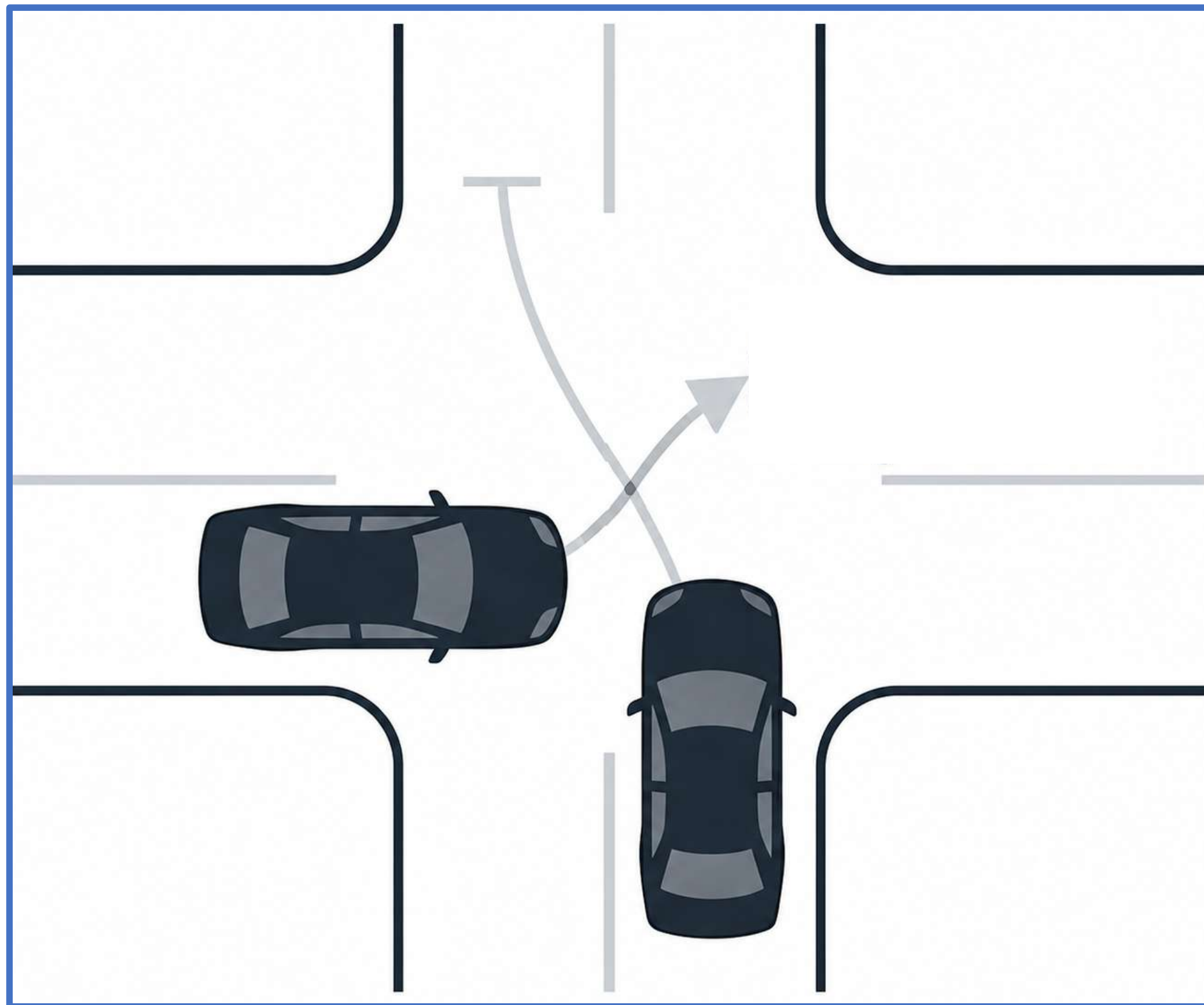


A. Naziri, S. Lambertenghi, M. D'Amorim. Misbehaviour Forecasting for Focused Autonomous Driving Systems Testing. Misbehaviour Forecasting for Focused Autonomous Driving Systems Testing. In Proceedings of the 48th International Conference on Software Engineering, 2026

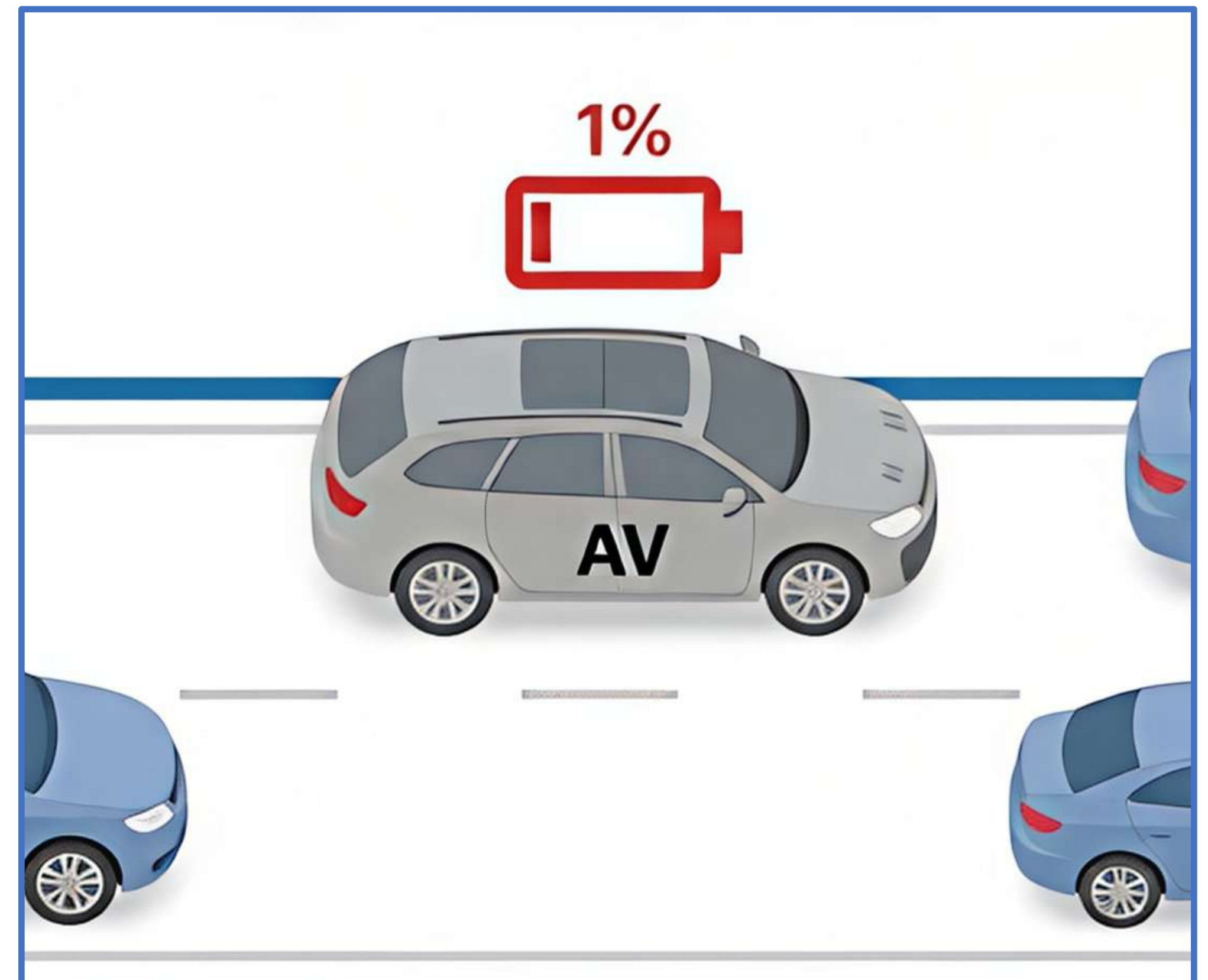
Other applications

Possible Field Directions

Safety-critical scenarios

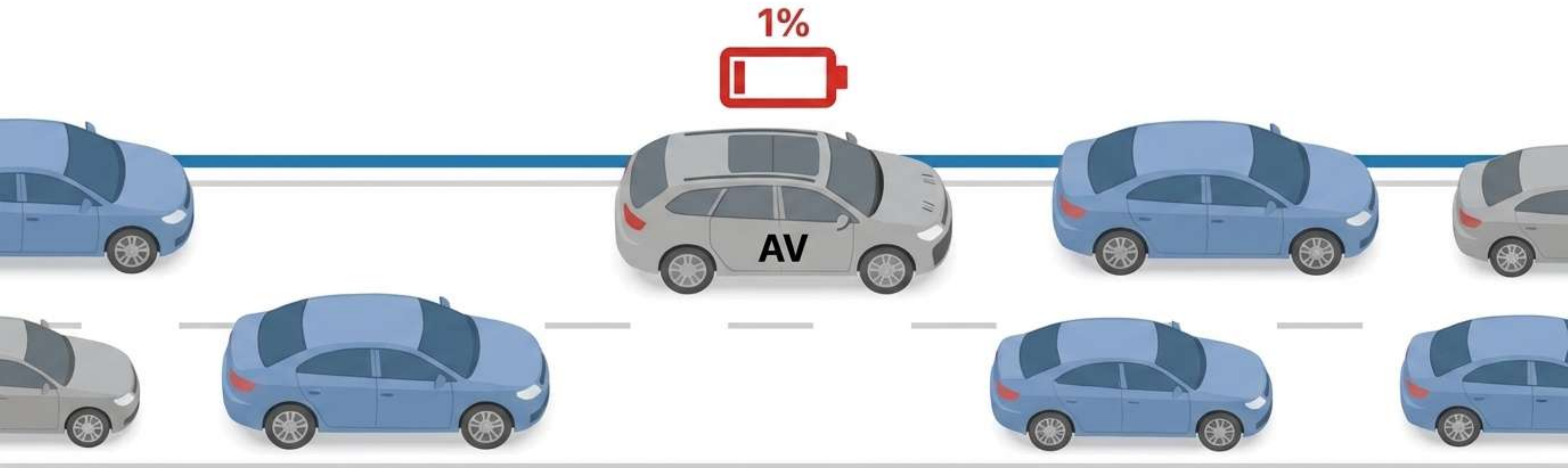


Energy-critical scenarios



Energy Critical Scenarios

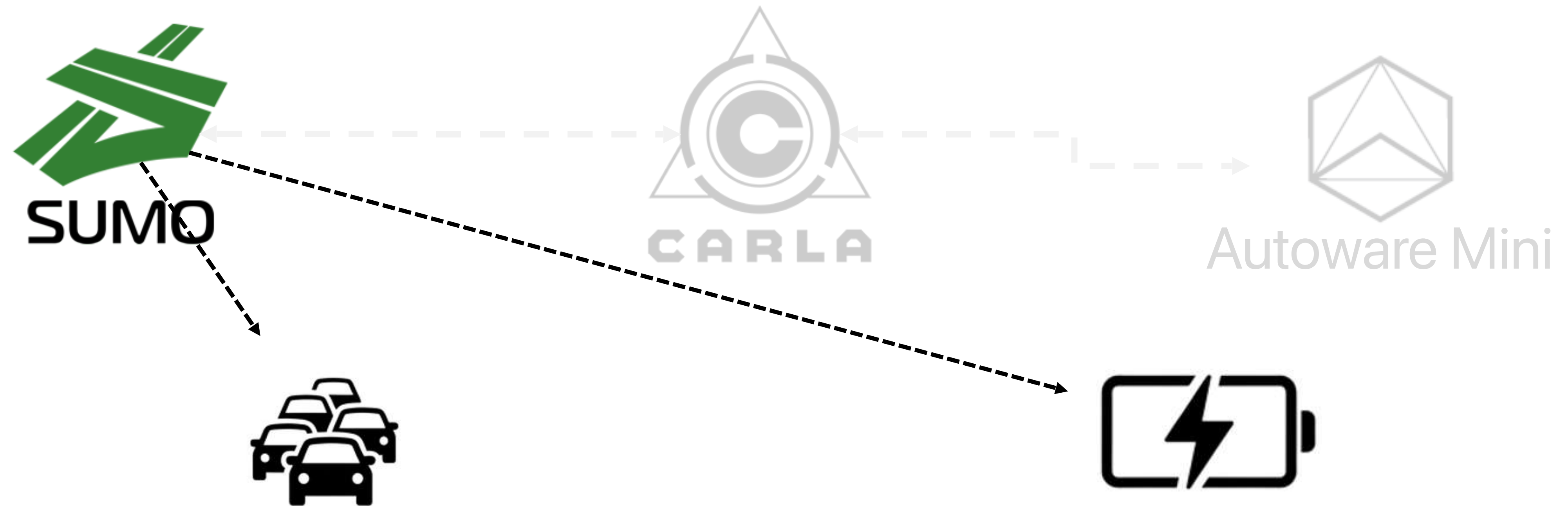
Non-functional. Unexplored field. Difficult to reproduce in reality.



Scenario Generator Tool



Scenario Generator Tool



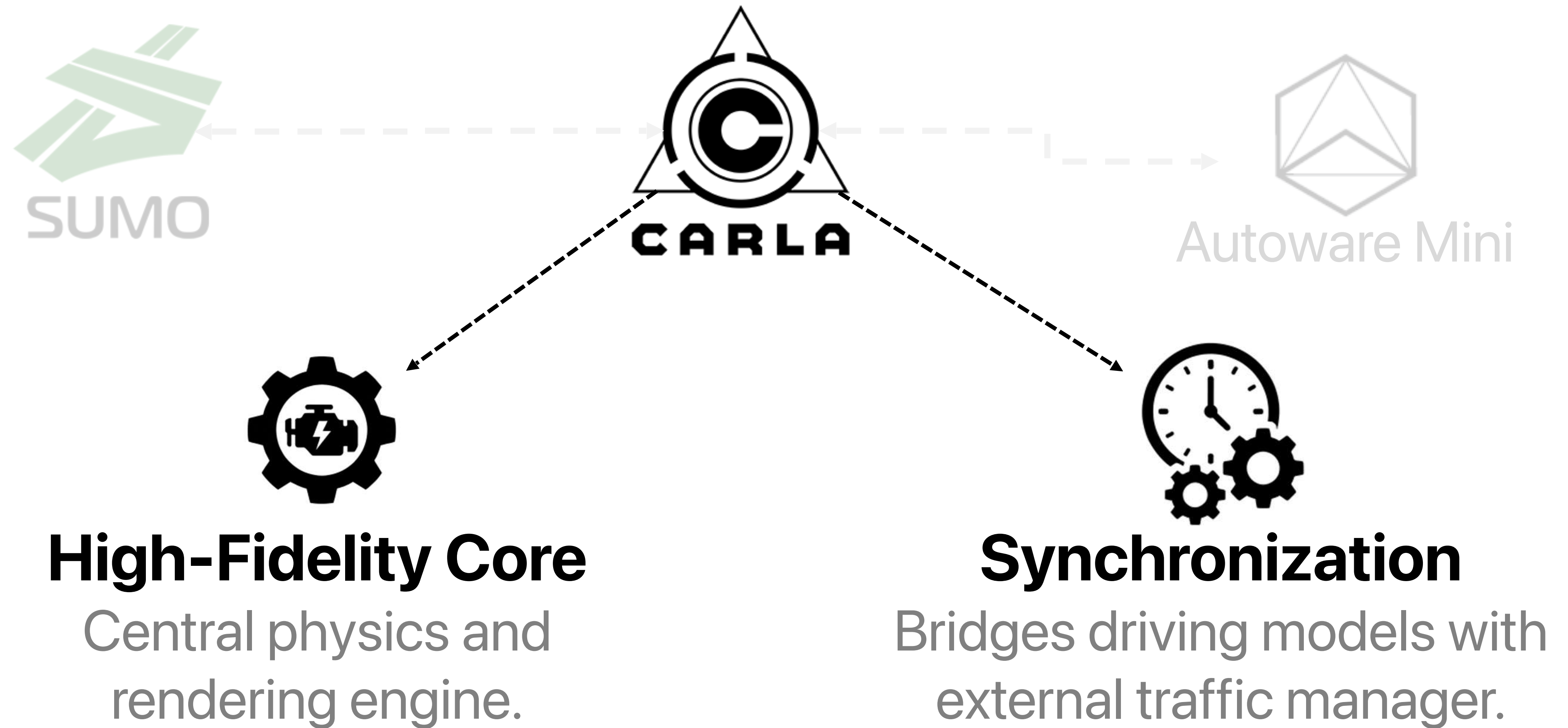
Traffic Manager

Controls NPC behavior,
routing and macro-level
Congestion.

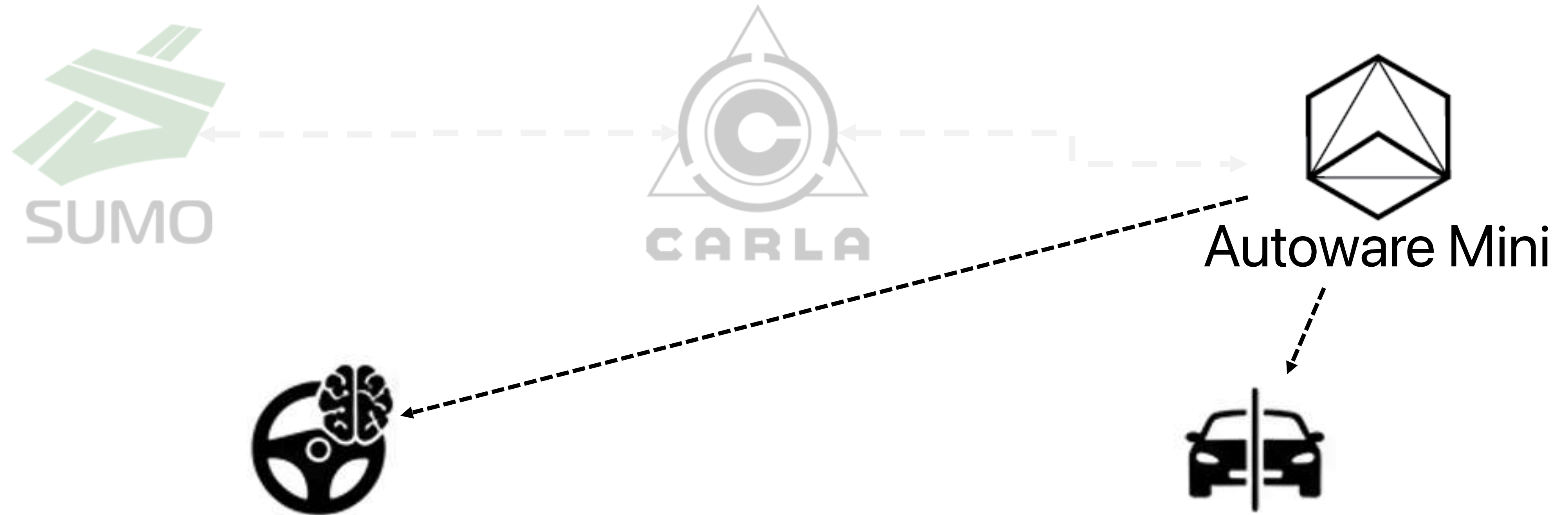
Electrical Manager

Manages the AV vehicle's
physical and electrical
characteristics.

Scenario Generator Tool



Scenario Generator Tool



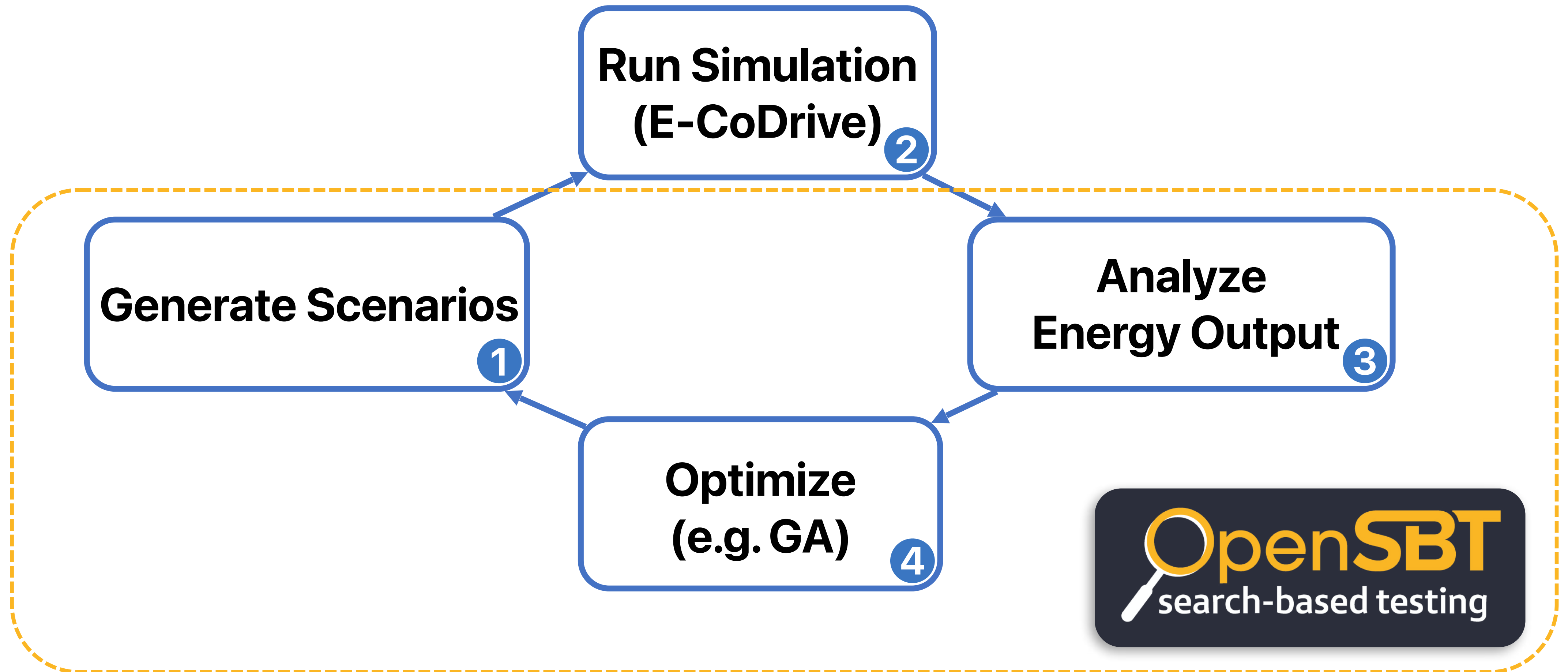
Autonomous Driving Software

Handles routing, perception,
and control.

Reality Translation

Can be used together with the
Real ADS software stack.

Search-Based Testing Application



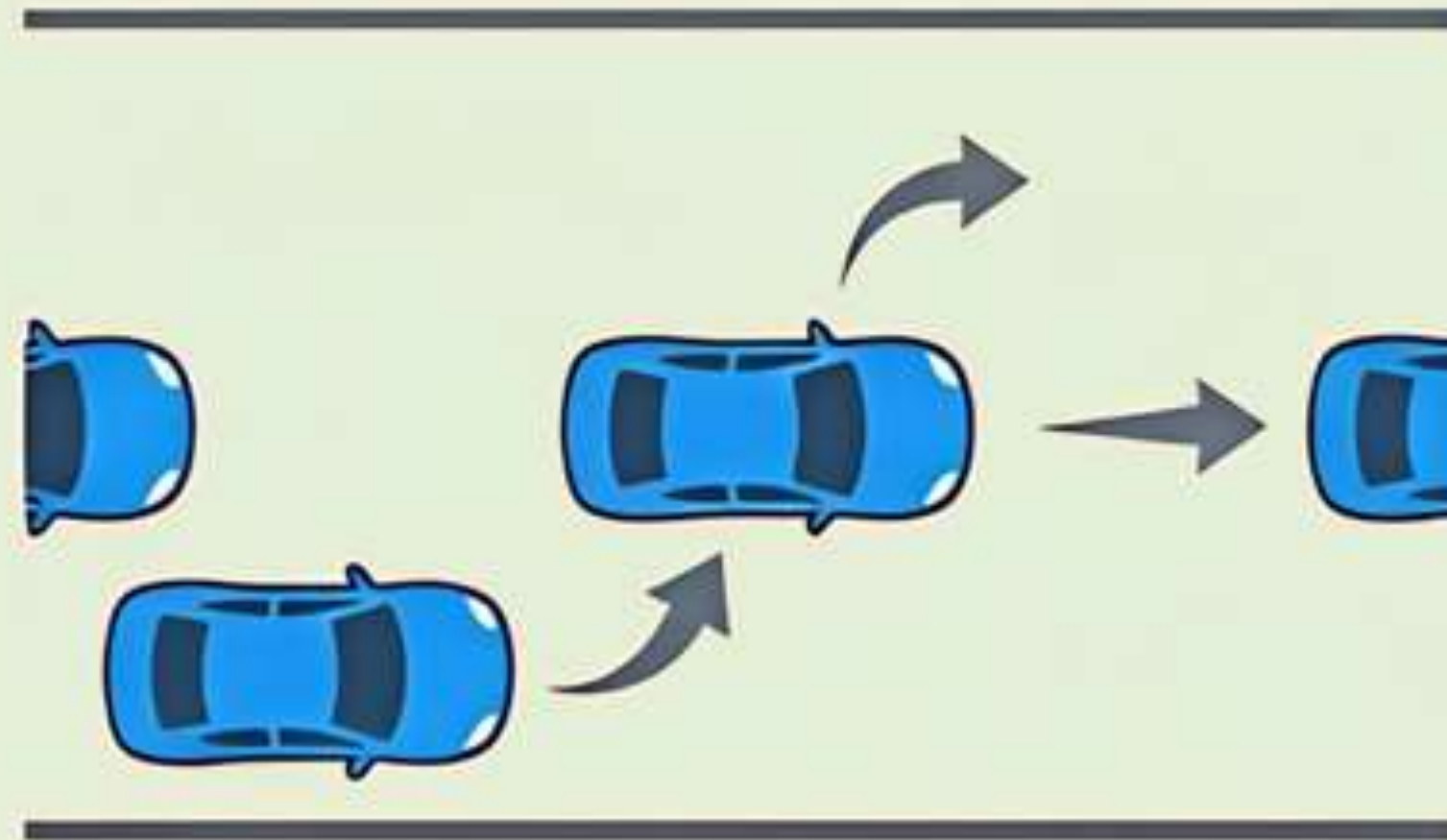
Spectrum of Traffic Scenario

Baseline



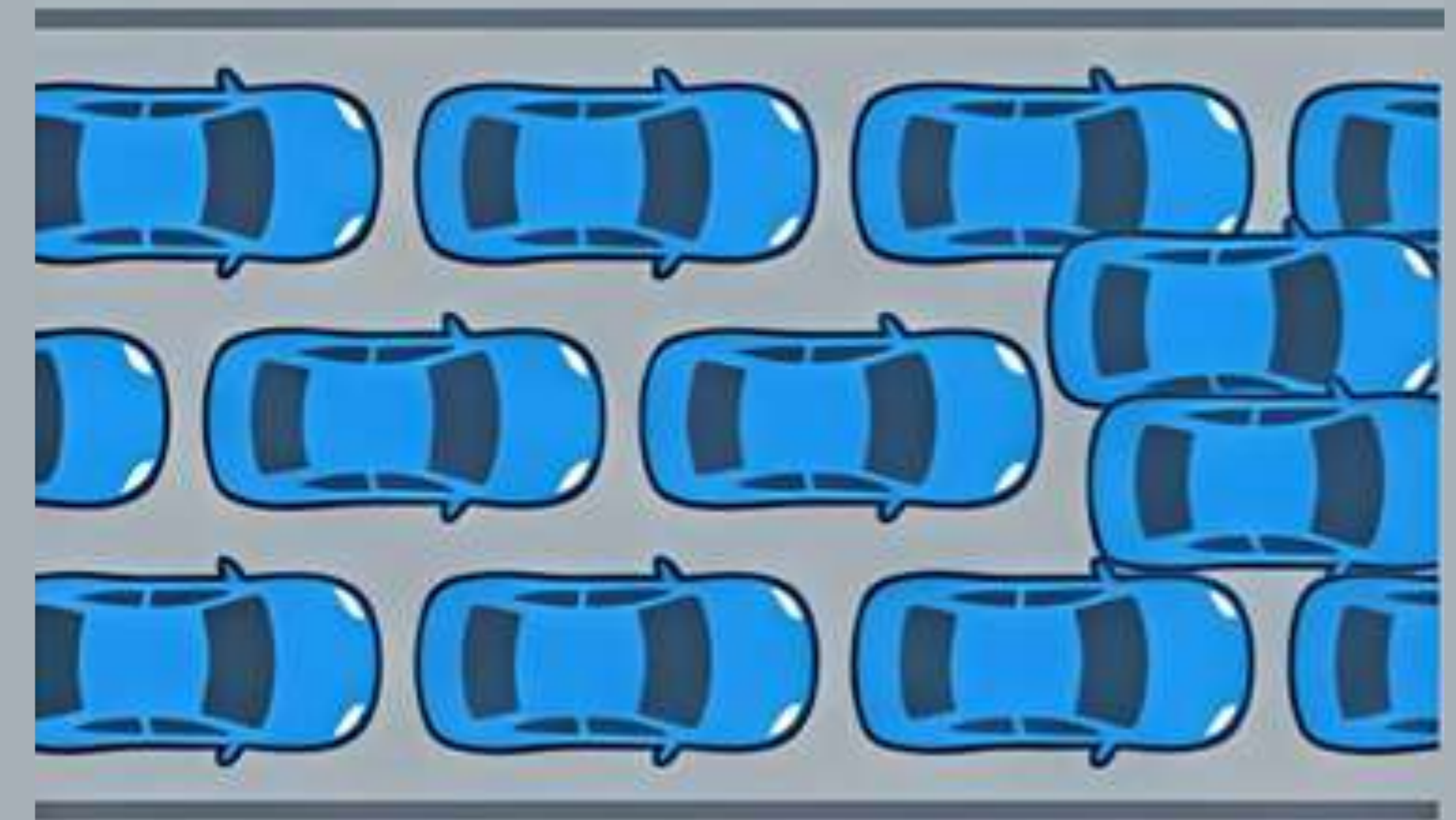
Low battery consumption.
Used as reference.

Sweet spot



High battery consumption due to
frequent acceleration/deceleration.
Not easy to find.

Trivial solution



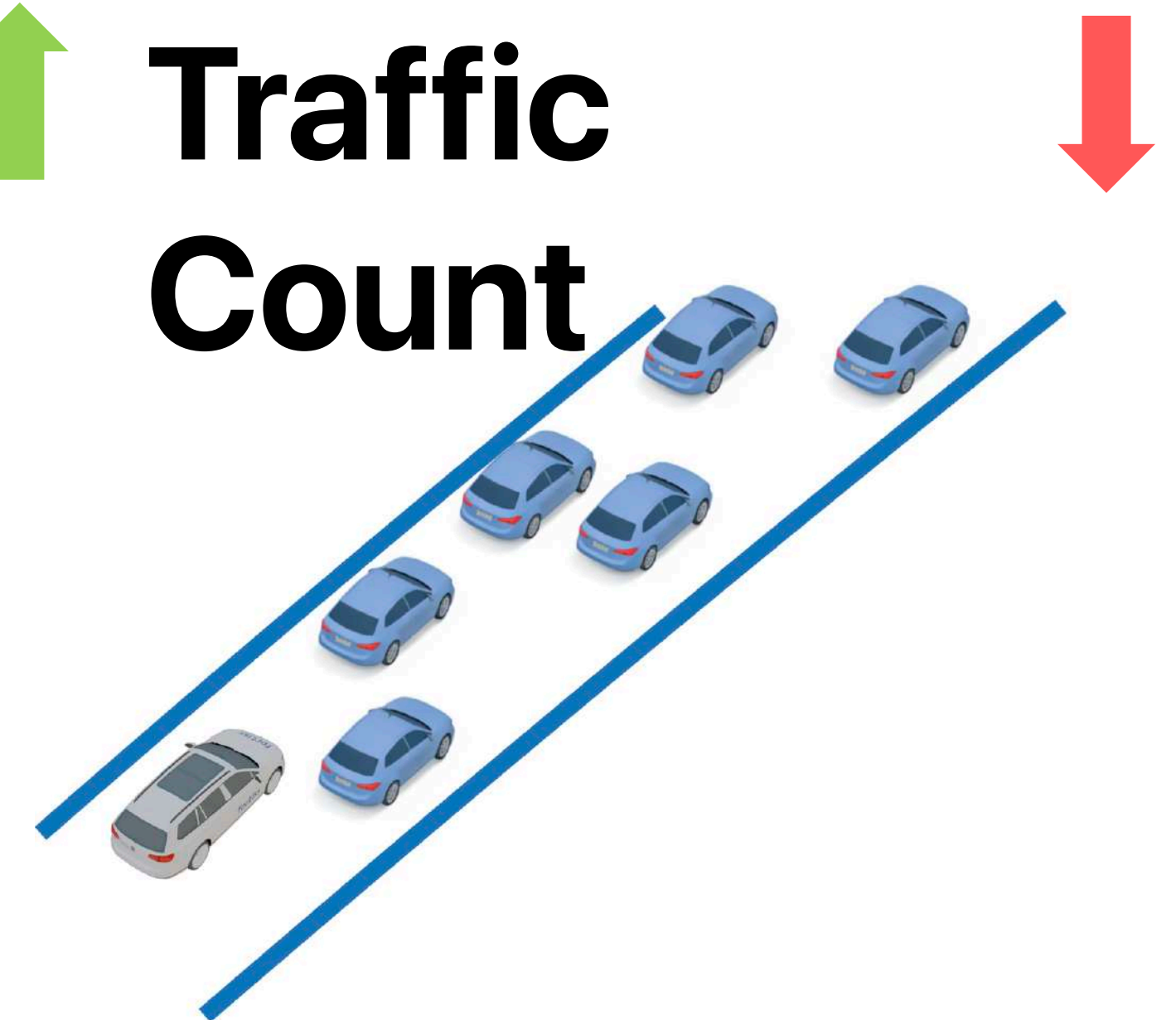
An easy-to-implement solution.
It does not promptly show the
impact of traffic.

Fitness Function Design

AEV speed ↑



Energy Consumption ↑ **Traffic Count** ↓



Preliminary Results



Effectiveness



Efficiency

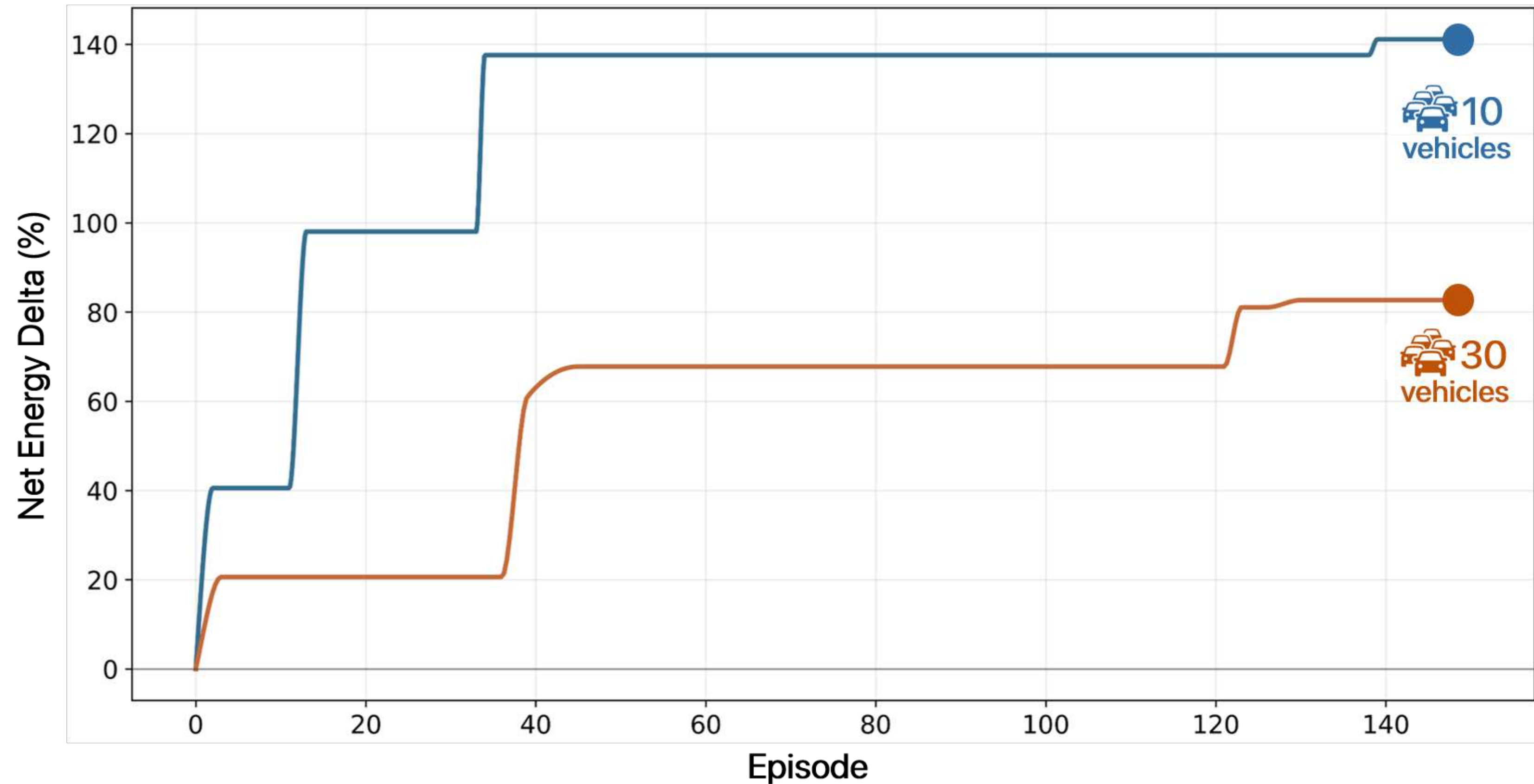


Diversity



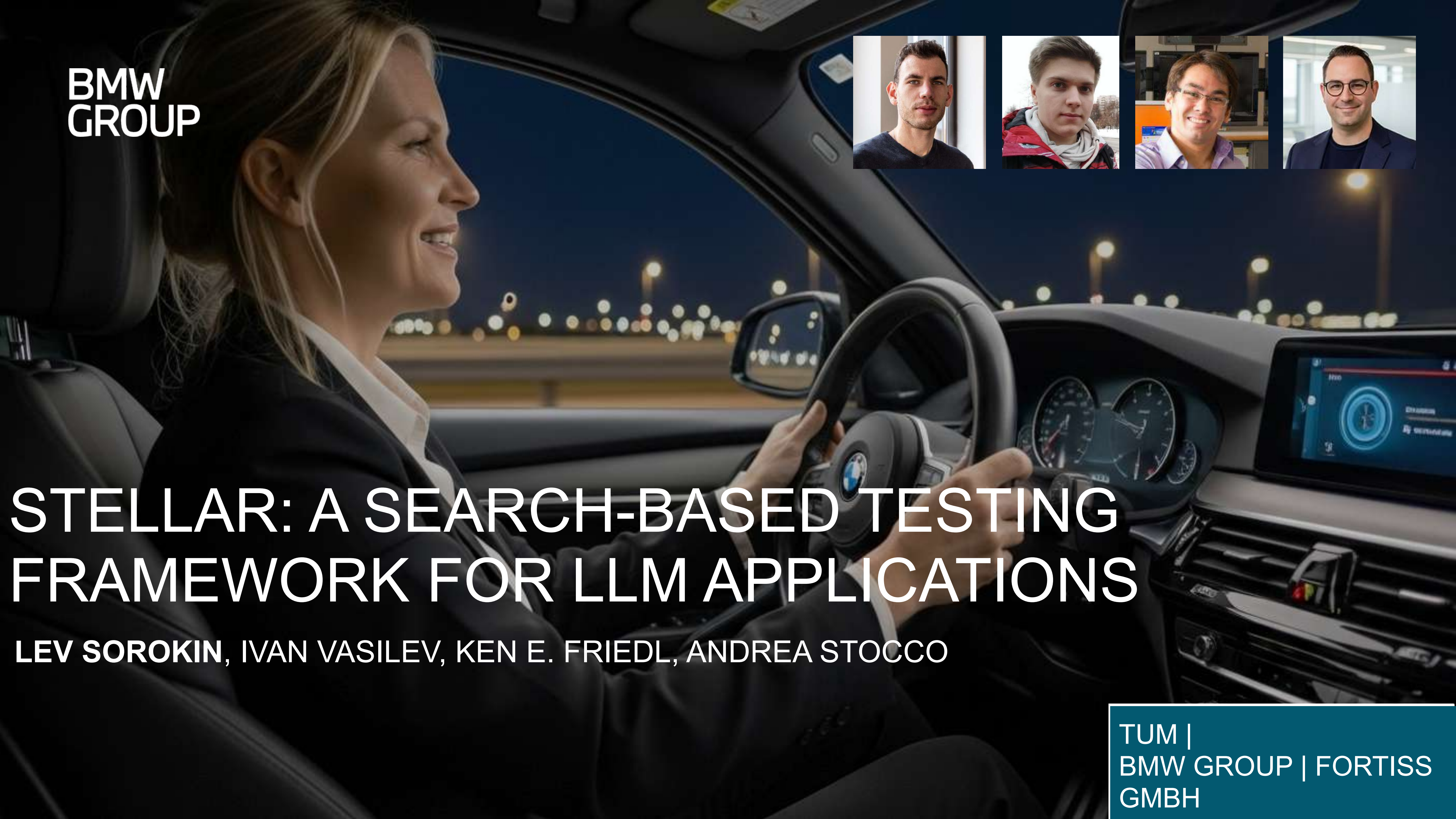
Transferability

Worst Energy Consumption over episodes



Conversational Agents in 1982





BMW
GROUP



STELLAR: A SEARCH-BASED TESTING FRAMEWORK FOR LLM APPLICATIONS

LEV SOROKIN, IVAN VASILEV, KEN E. FRIEDL, ANDREA STOCCO

TUM |
BMW GROUP | FORTISS
GMBH

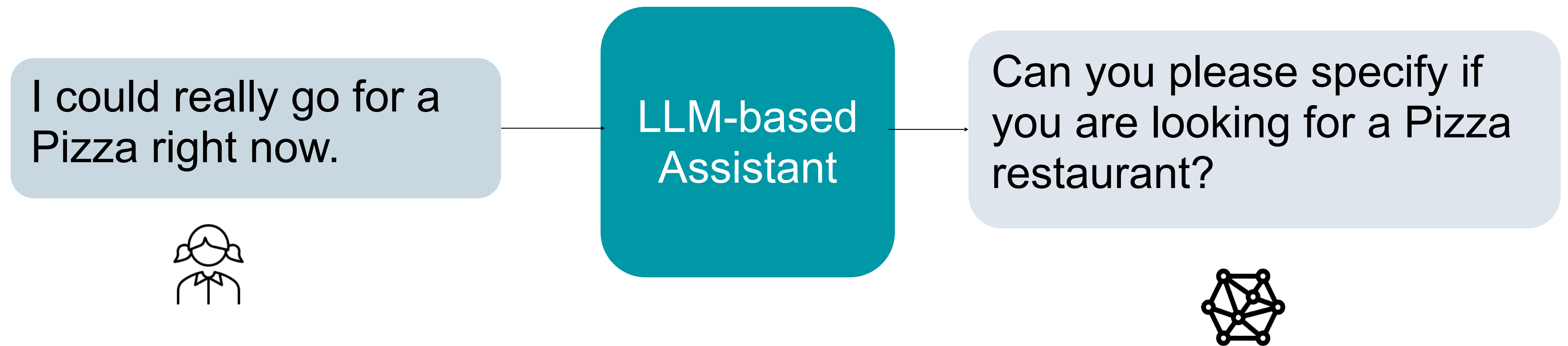
Motivation



Example Malicious User Input



Example Implicit Request



PROBLEM

Test space huge: defined by **expressional** and **robustness** features

- Formality
- Slang
- Implicitness
- Anthropomorphism
- Misspelling
- Fillers...



I need Pizza now!



I would like to it Pizza now.



Show me a Pizza spot, please.



I would like to, ehm, get Pizza.



Do you know a good Pizza place?

TEST REPRESENTATION

Content

venue: *[bar, restaurant, pizzeria,,...]*

price: *[\$, \$\$, \$\$\$]*

rating: *[3, 3.5, 4, 4.5, 5]*

...

Style

politeness: *[„polite“, „rude“]*

implicit: *[„implicit“, “slightly implicit“]*

...

Perturbation

missing words: *[0%, 10%, ..., 50%]*

use_fillers: *[yes, no]*

...

Domain Specification

venue: *car_repair*
→ food_type: *None*

venue: *hospital*
→ food_type: *None*

venue: *bar*
→ fuel_price: *None*

...

Feature Constraints

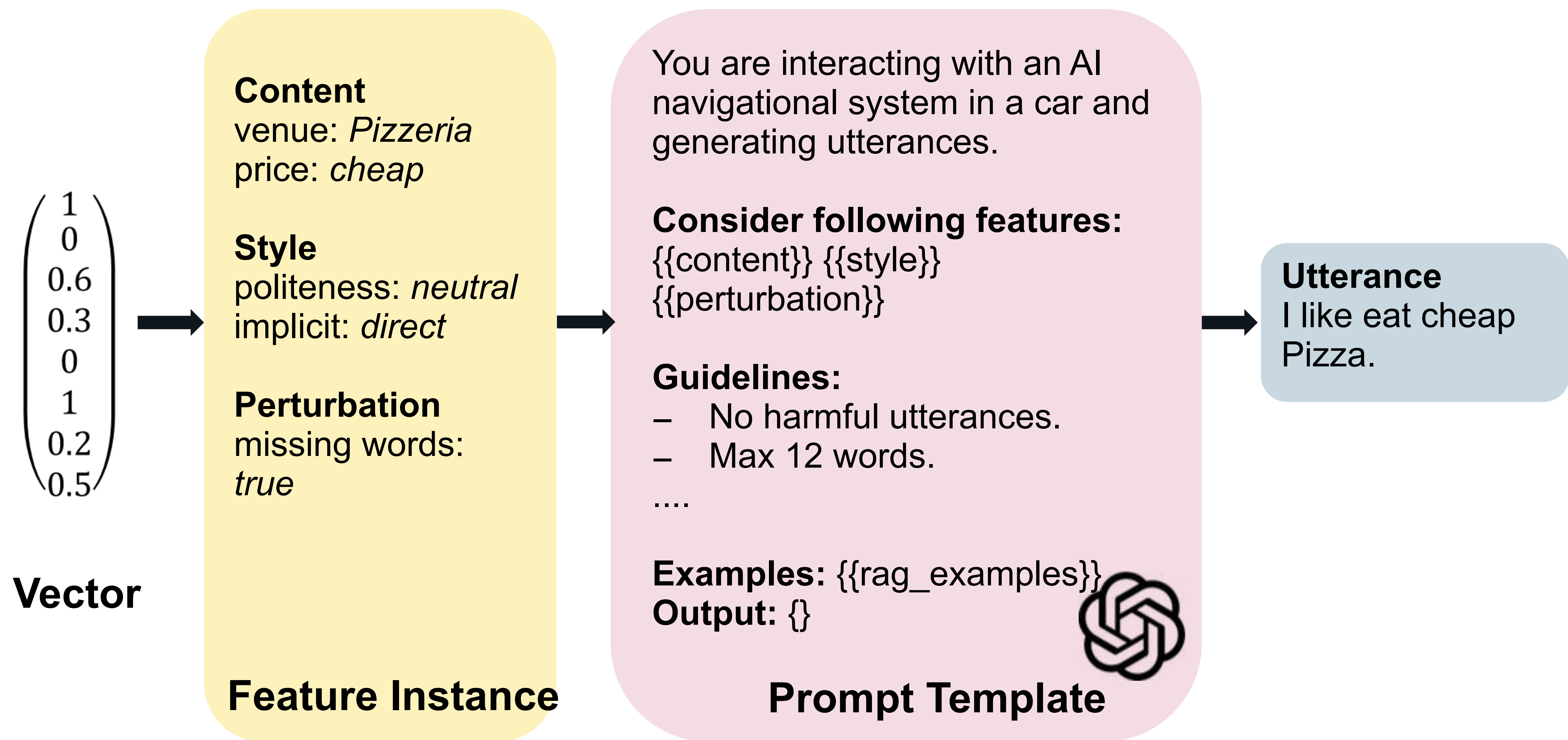
Content = $(0, 0.3, 0.5, 0.1)$

Style = $(0, 0.3, 0.6, 0.3)$

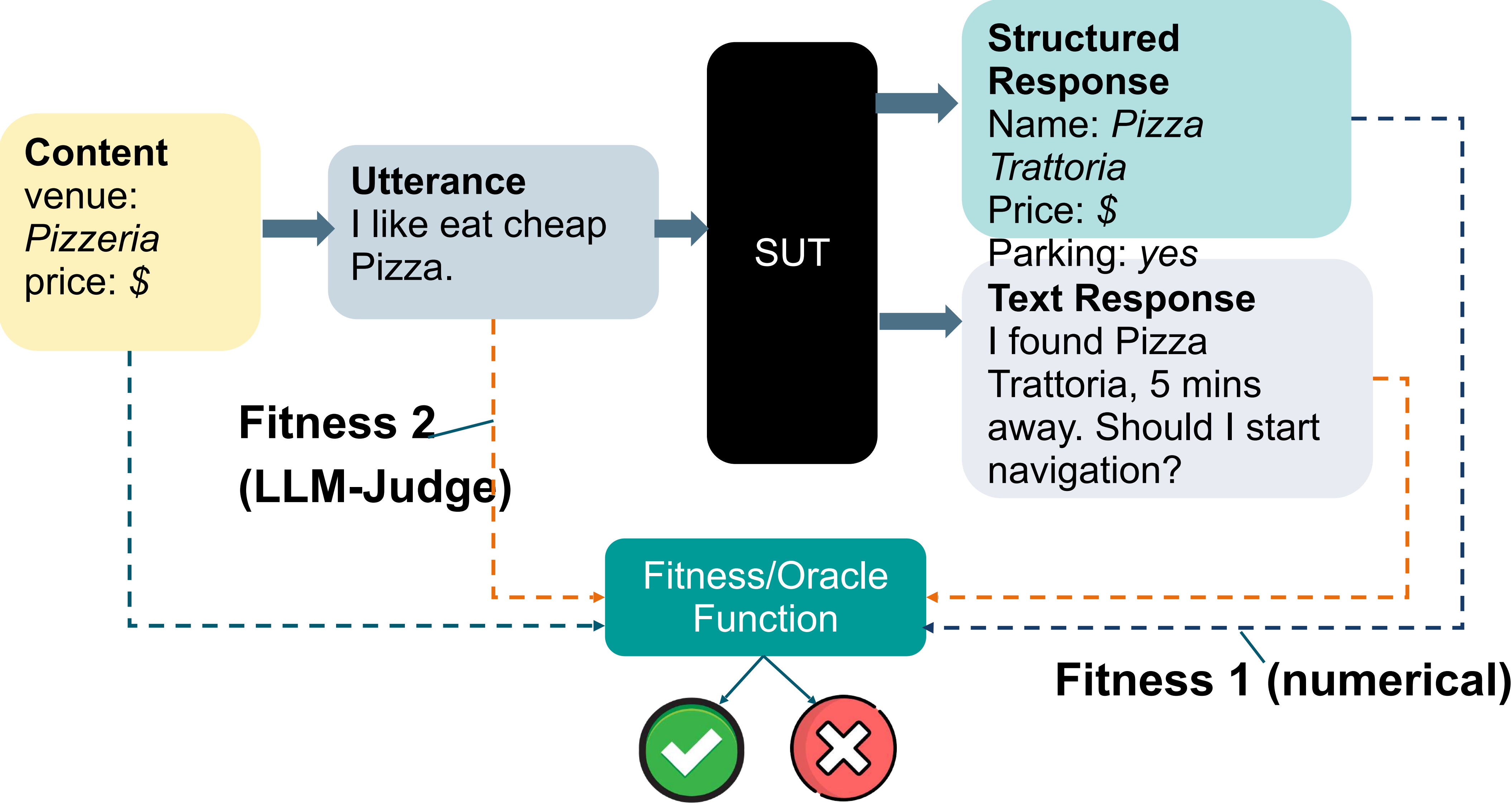
Perturb = $(0, 0.3, 0.6, 0.3)$

Numerical Encoding

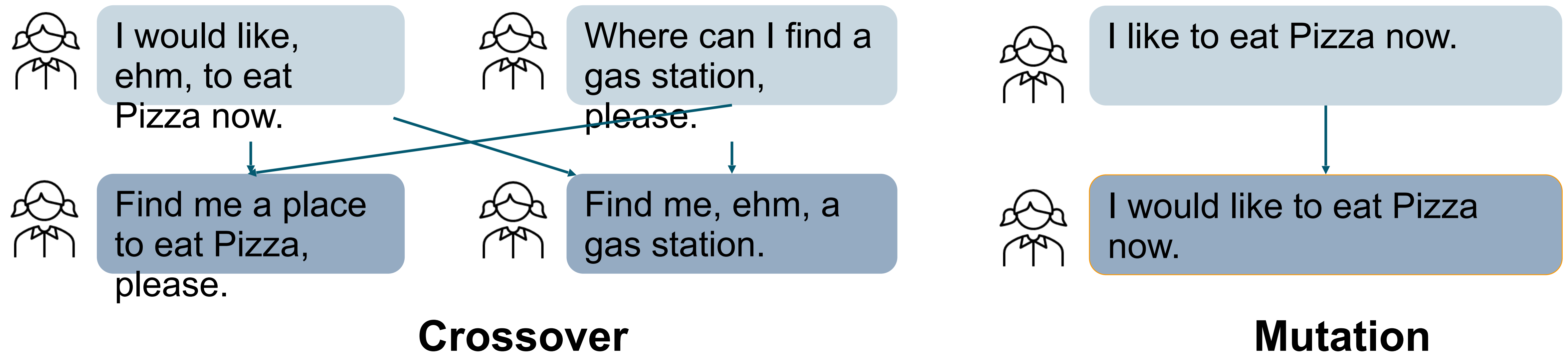
STEP 2 | TEST GENERATION



STEP 3 | TEST EXECUTION AND EVALUATION



STEP 4 | TEST OPTIMIZATION



Optimization by **applying genetic operations** on **numerical** features

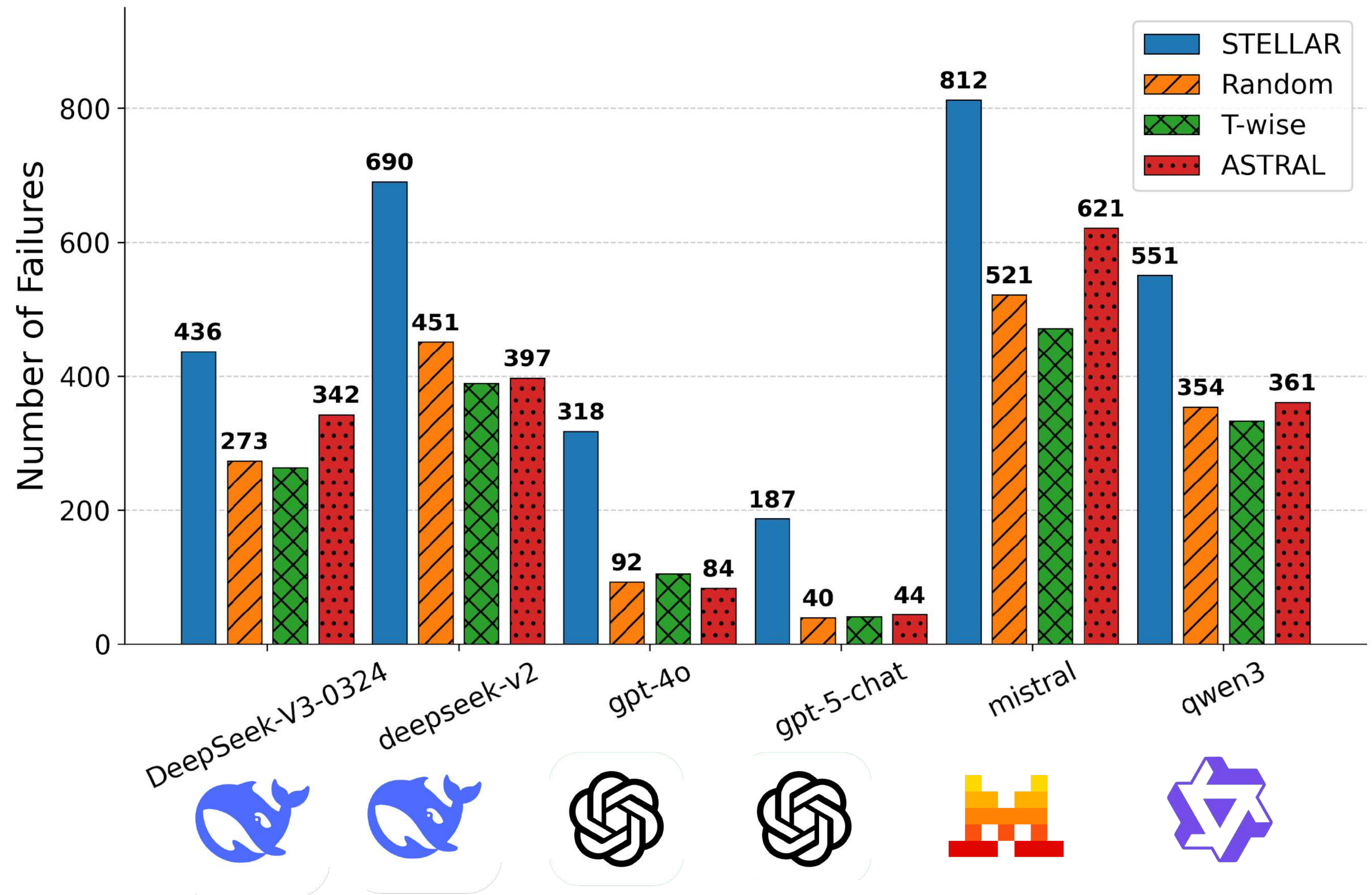
Past executions **guide** generation of new test inputs

EMPIRICAL EVALUATION

Generation of more than **200,000 tests**
(> 200 hours)

STELLAR exposes up to **4.3×** more failures than existing approaches

STELLAR increases **coverage** of failures for strong models



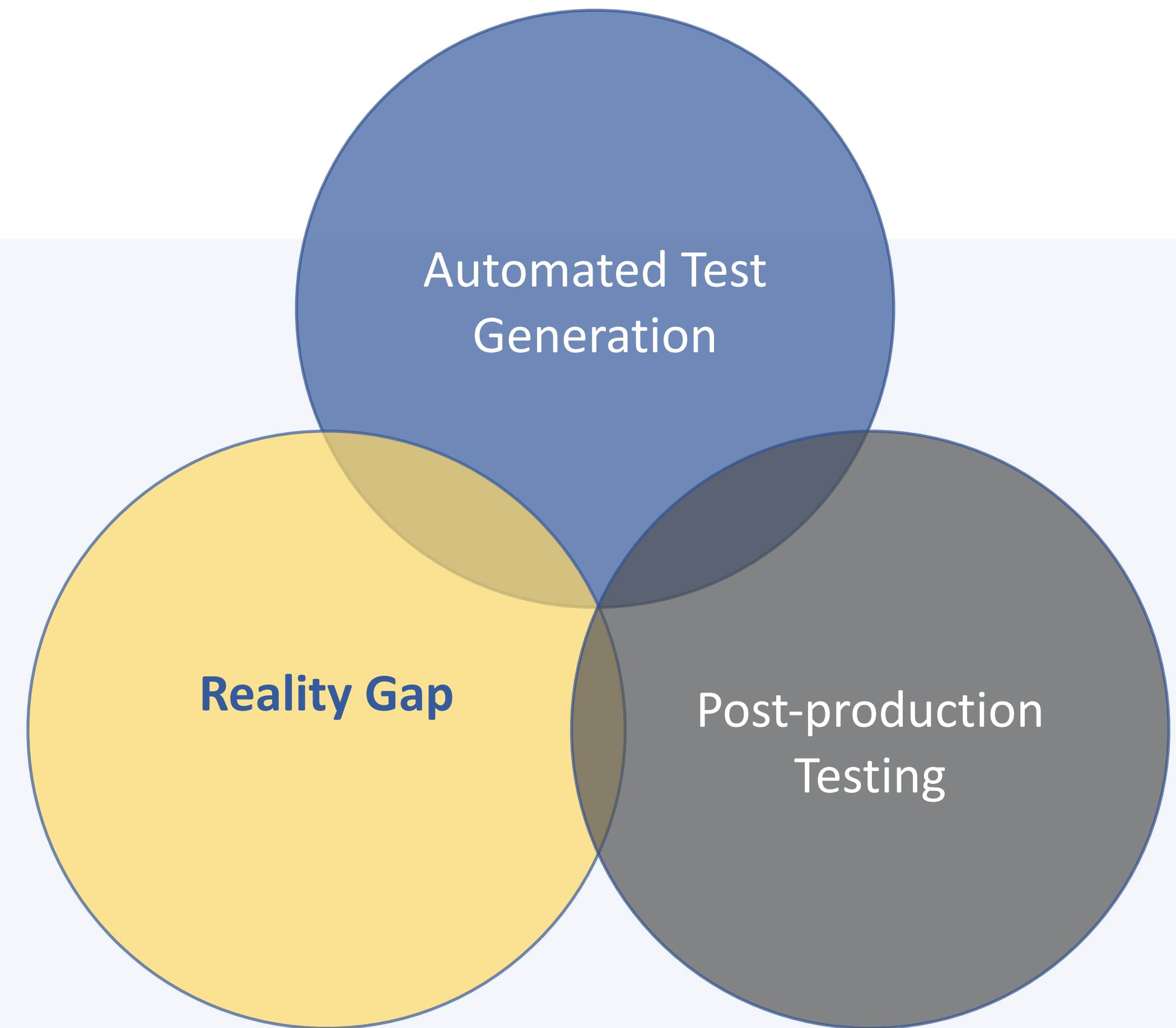
Takeaways

- ① **Test generation encompasses three objectives: validity, effectiveness, efficiency**
- ② **Multi-simulator agreement improves the validity of detected failures**
- ③ **Open challenges: more representative simulators, different fidelities, neural simulations, hybrid simulations (e.g., hardware in the loop), more accessible ADS for testing**

Automated Software Testing

Main research topics

- **Automated Test Generation**
How can we automatically generate complex scenario-based tests efficiently and effectively?
- **Reality Gap**
How can we measure and mitigate the reality gap between simulation and in-field testing?
- **Post-production Testing**
How to ensure a high dependability of deep neural network driven-cyber-physical systems (CPS) in production?



Automated Driving System (ADS) testing

Real-world testing

⊕ Evaluate ADS realistically

⊖ High cost
High risk
Constrained



Automated Driving System (ADS) testing

Real-world testing

⊕ Evaluate ADS realistically

⊖ High cost
High risk
Constrained

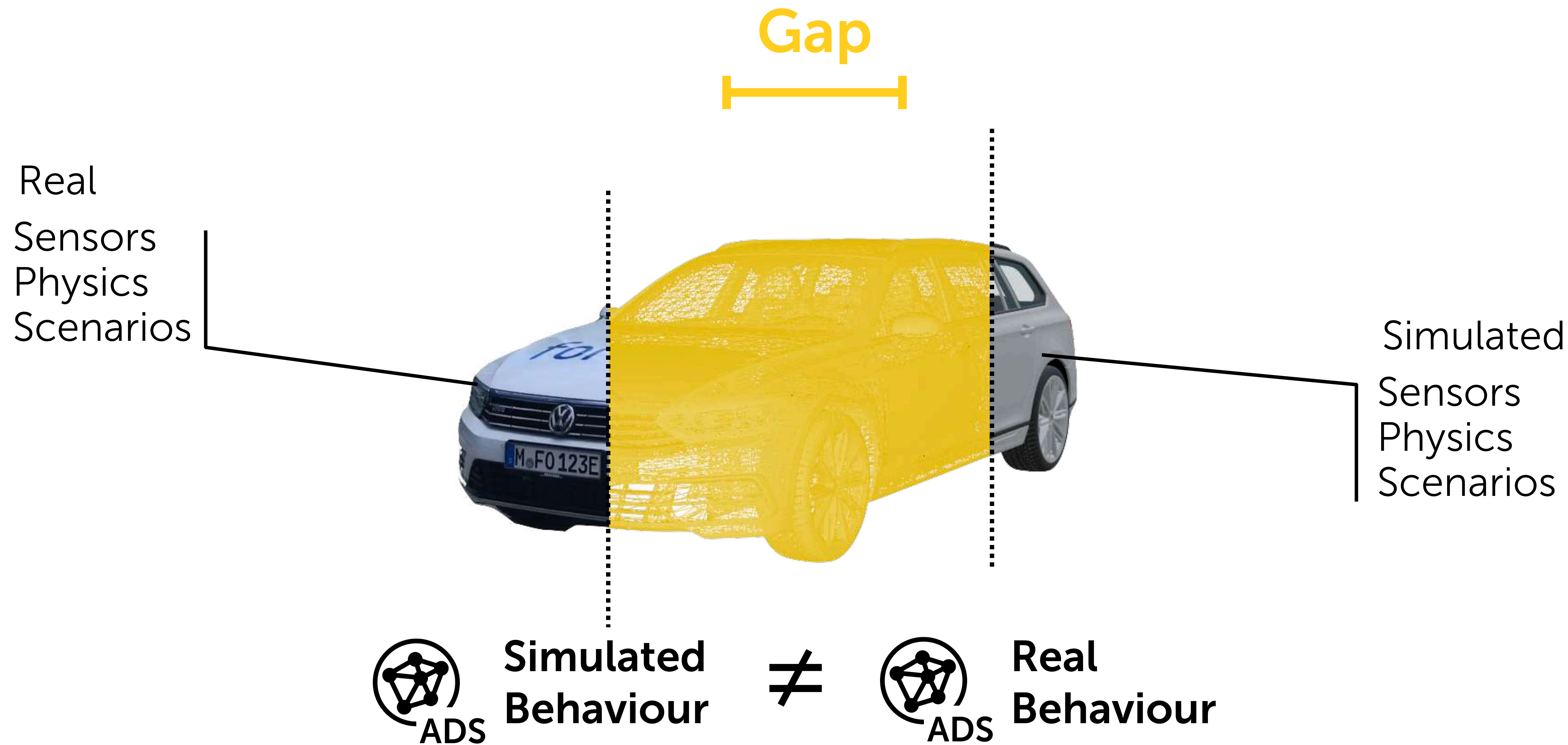


Simulation testing

⊕ Cheaper
Safer
More versatile

⊖ Approximation of reality

Reality Gap



Reality Gap

Simulated

Real



GAP

Sensors

Physics

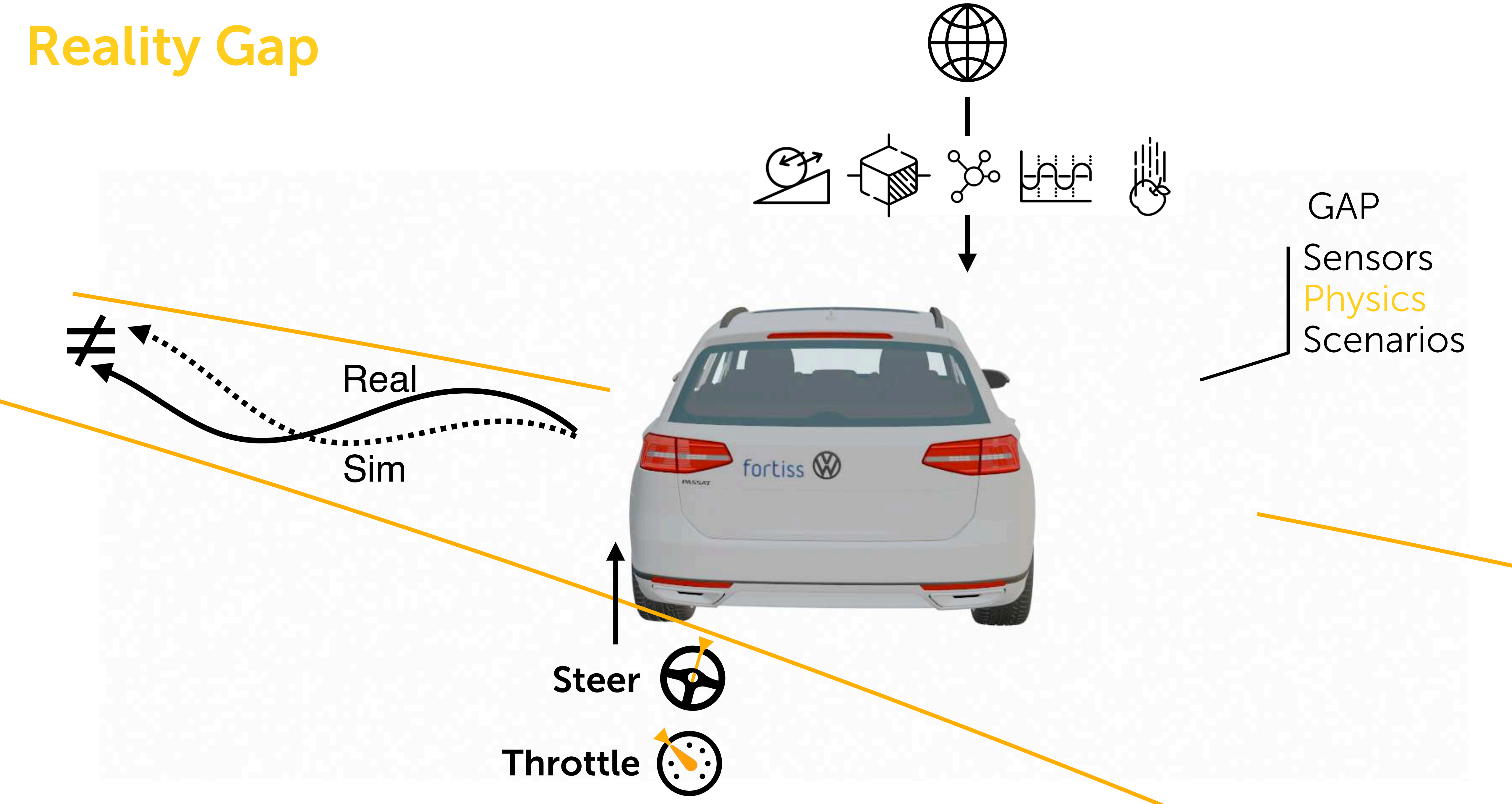
Scenarios



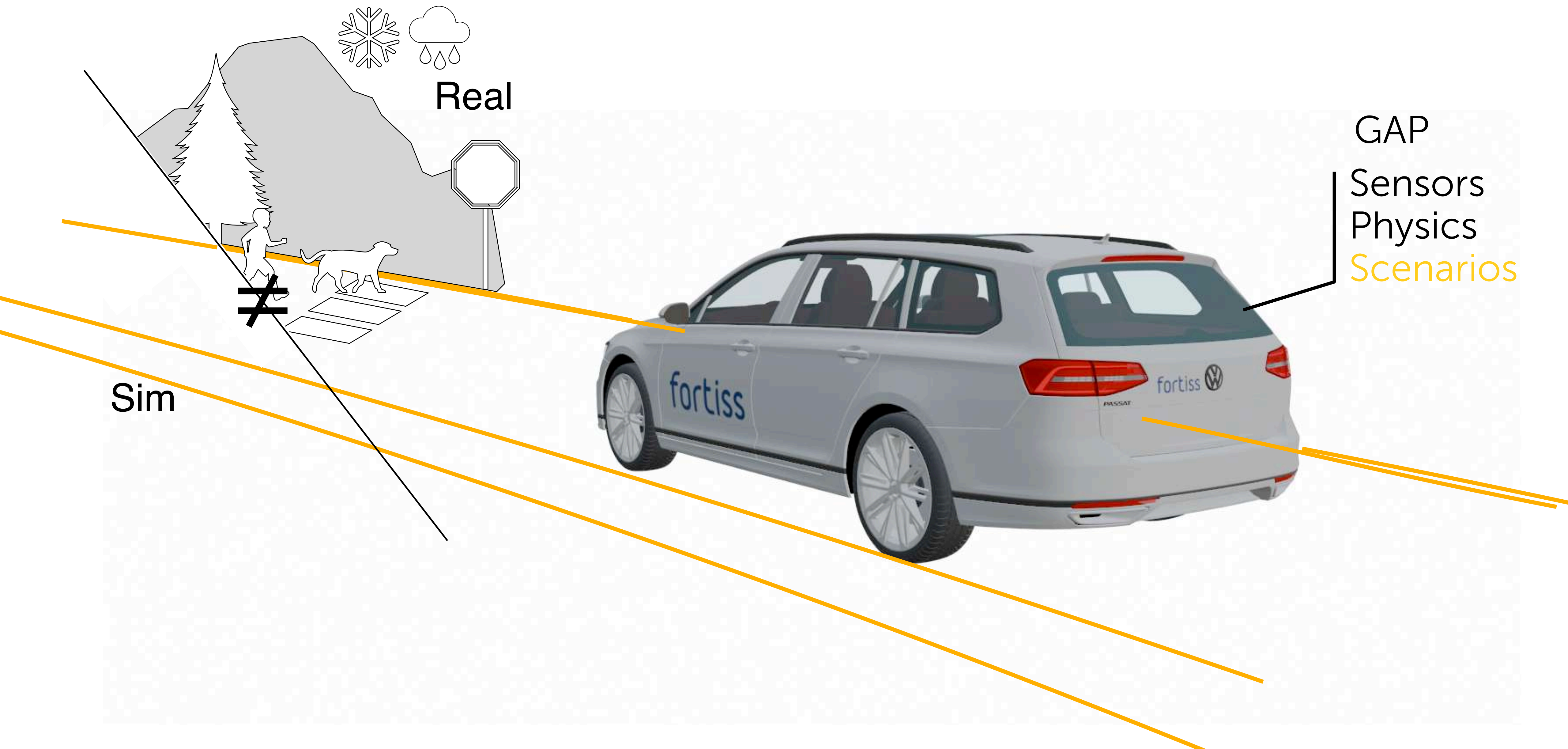
≠



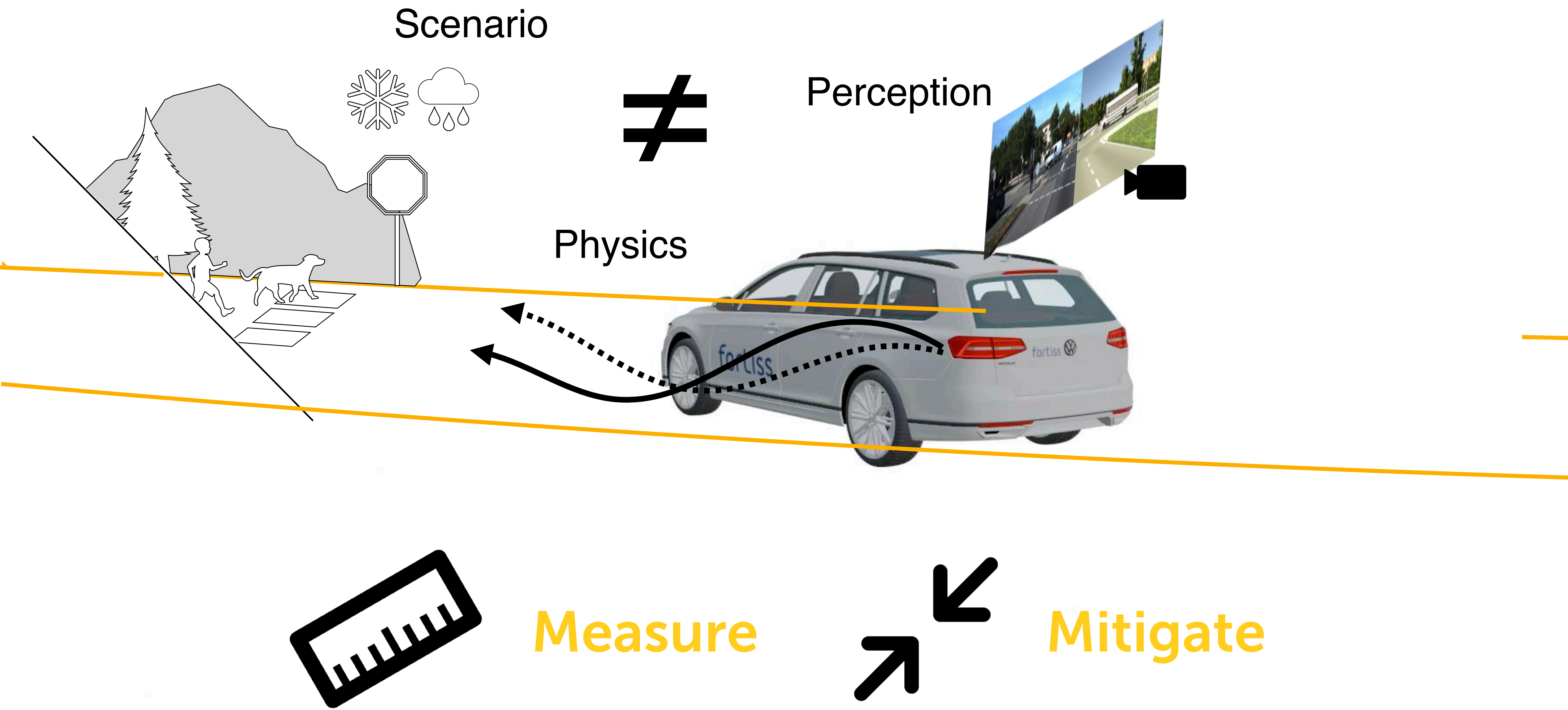
Reality Gap



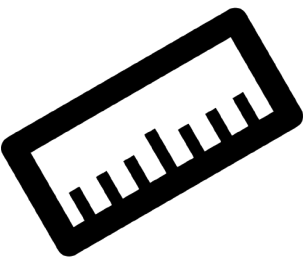
Reality Gap



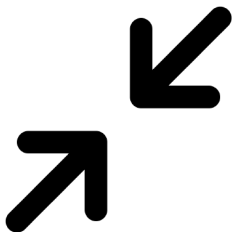
Our Goal



Our techniques

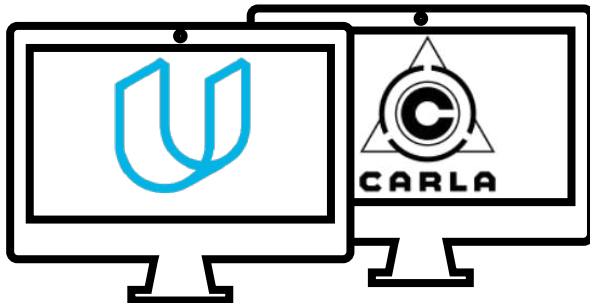


Measure



Mitigate

Generic
simulators



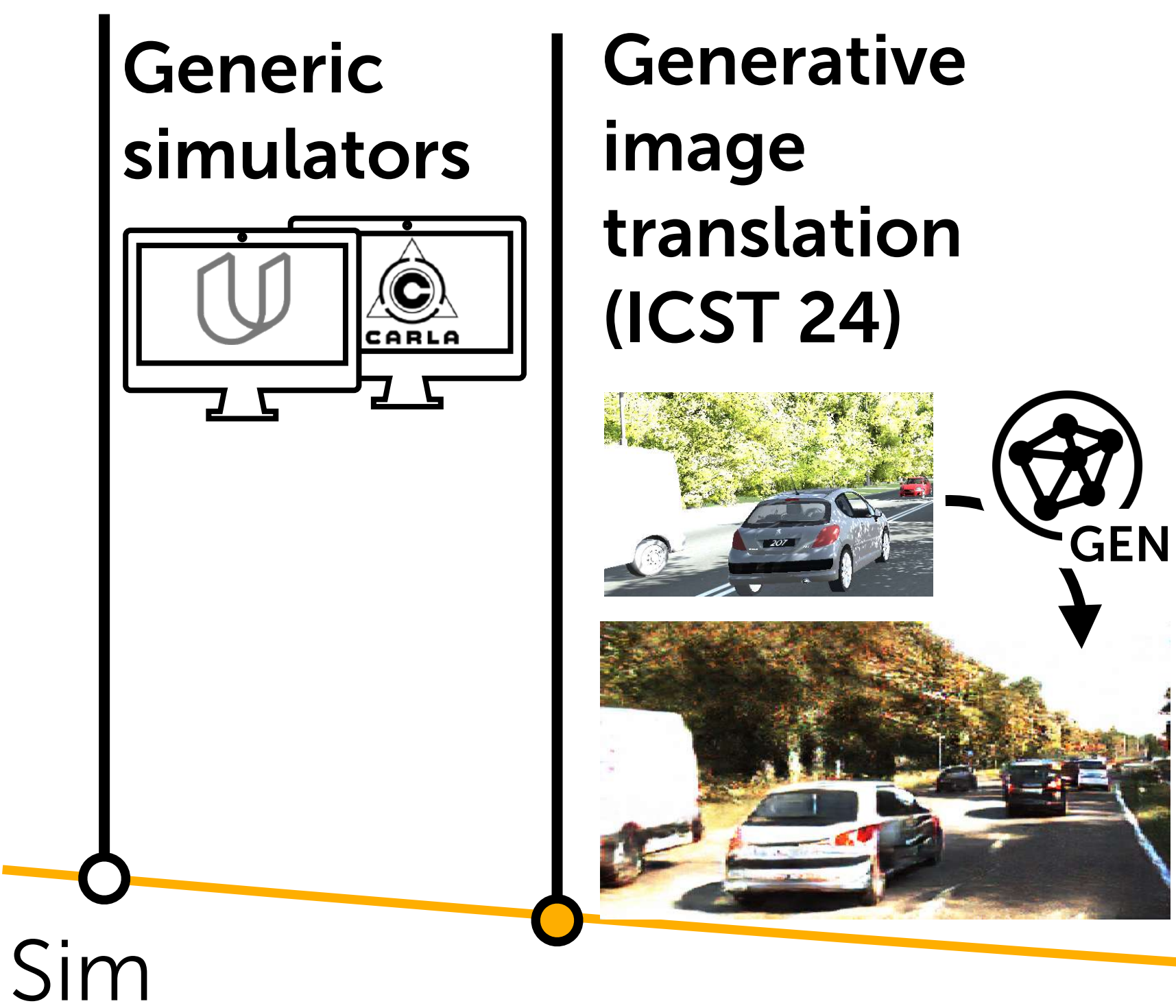
Sim



Real



Reality Gap Mitigation



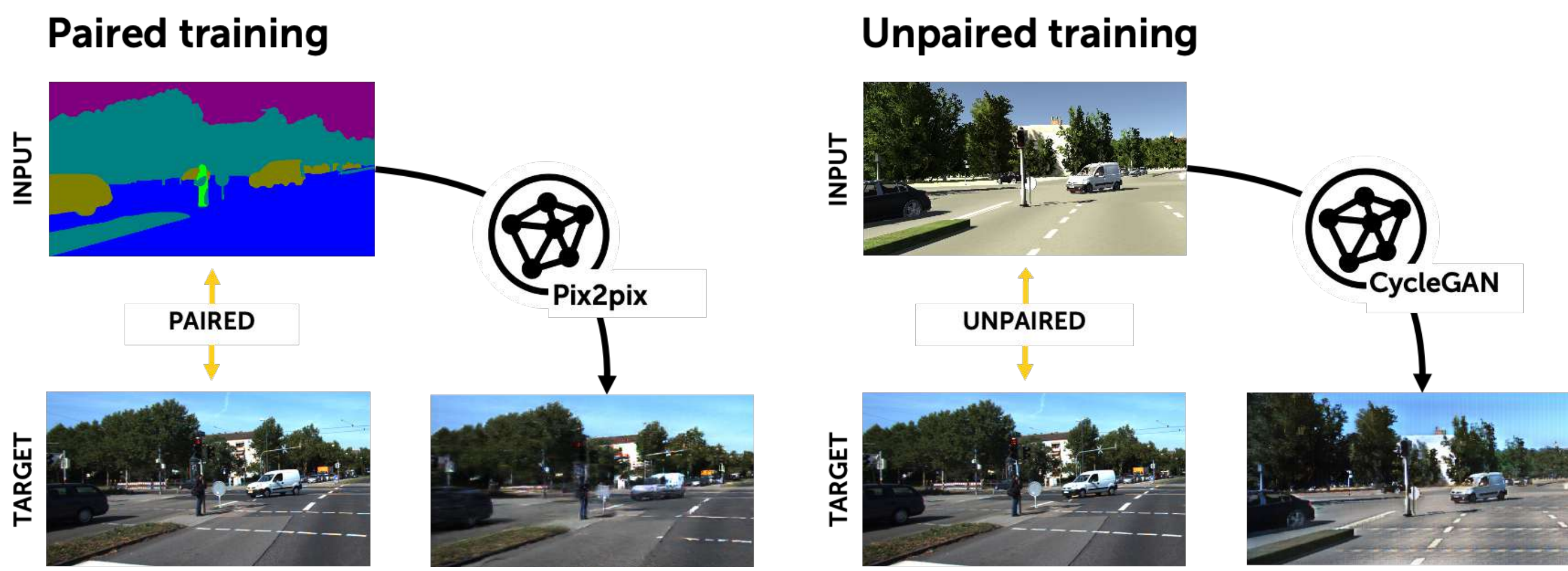
Sensors
Physics
Scenarios

Assessing Quality Metrics for Neural Reality Gap
Input Mitigation in Autonomous Driving Testing

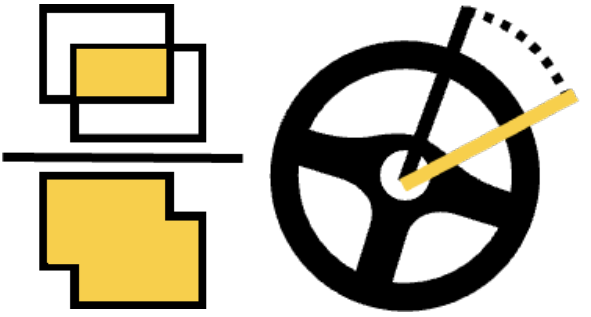
Stefano Carlo Lambertenghi
Technical University of Munich, fortiss GmbH
Munich, Germany
stefanocarlo.lambertenghi@tum.de, lambertenghi@fortiss.org

Andrea Stocco
Technical University of Munich, fortiss GmbH
Munich, Germany
andrea.stocco@tum.de, stocco@fortiss.org

ICST 24



IoU & MSE



VS

3
10



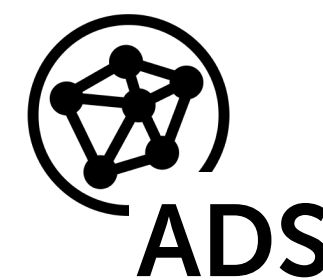
Distribution Level Metrics

Single Image Metrics

Perception Reality Gap

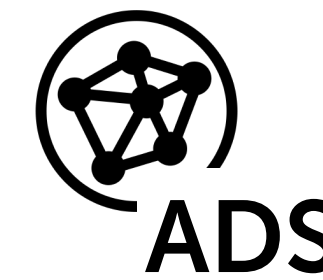
Difference between simulated and real input images

Simulation



Simulated
Behaviour

$\neq$

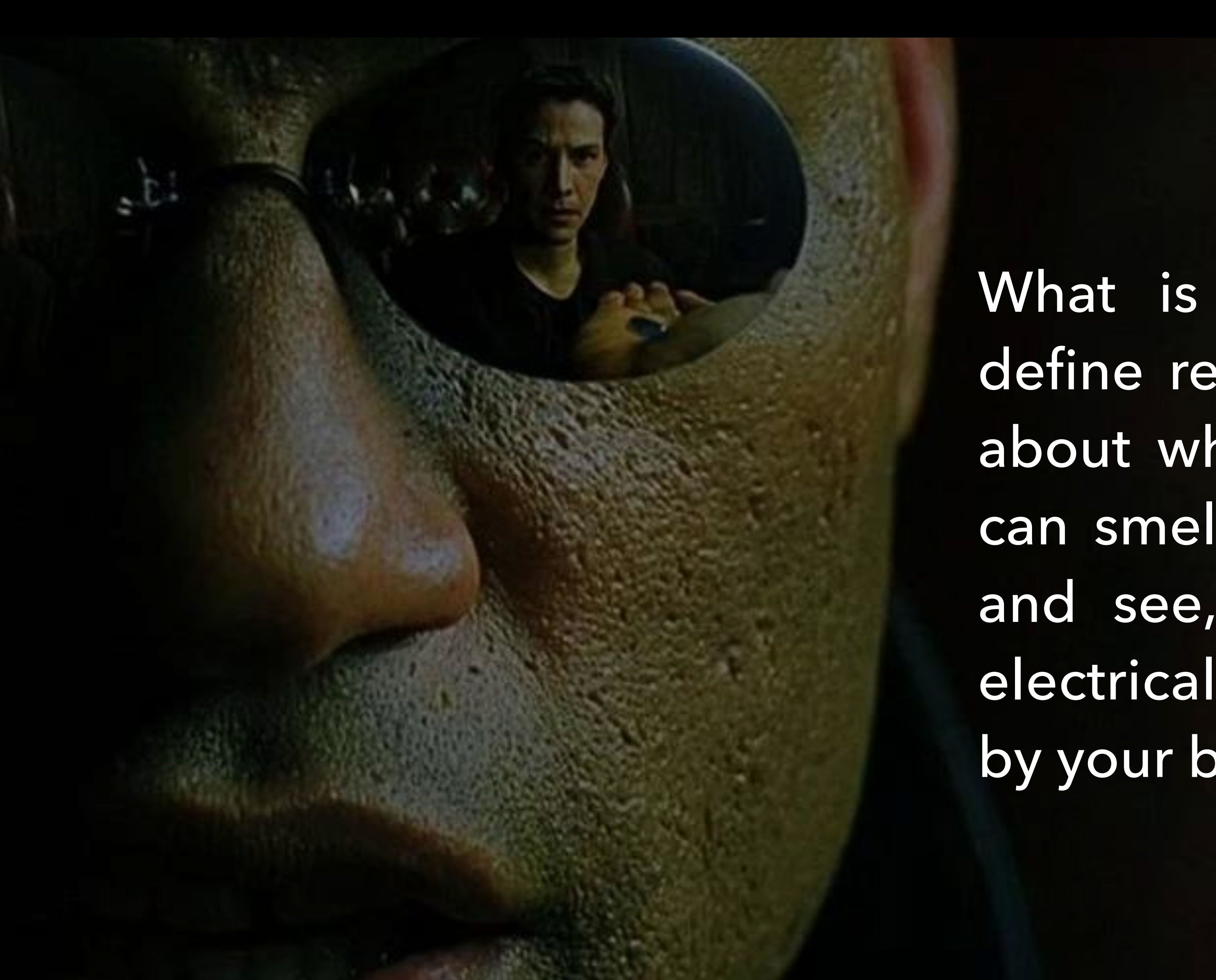


Real
Behaviour

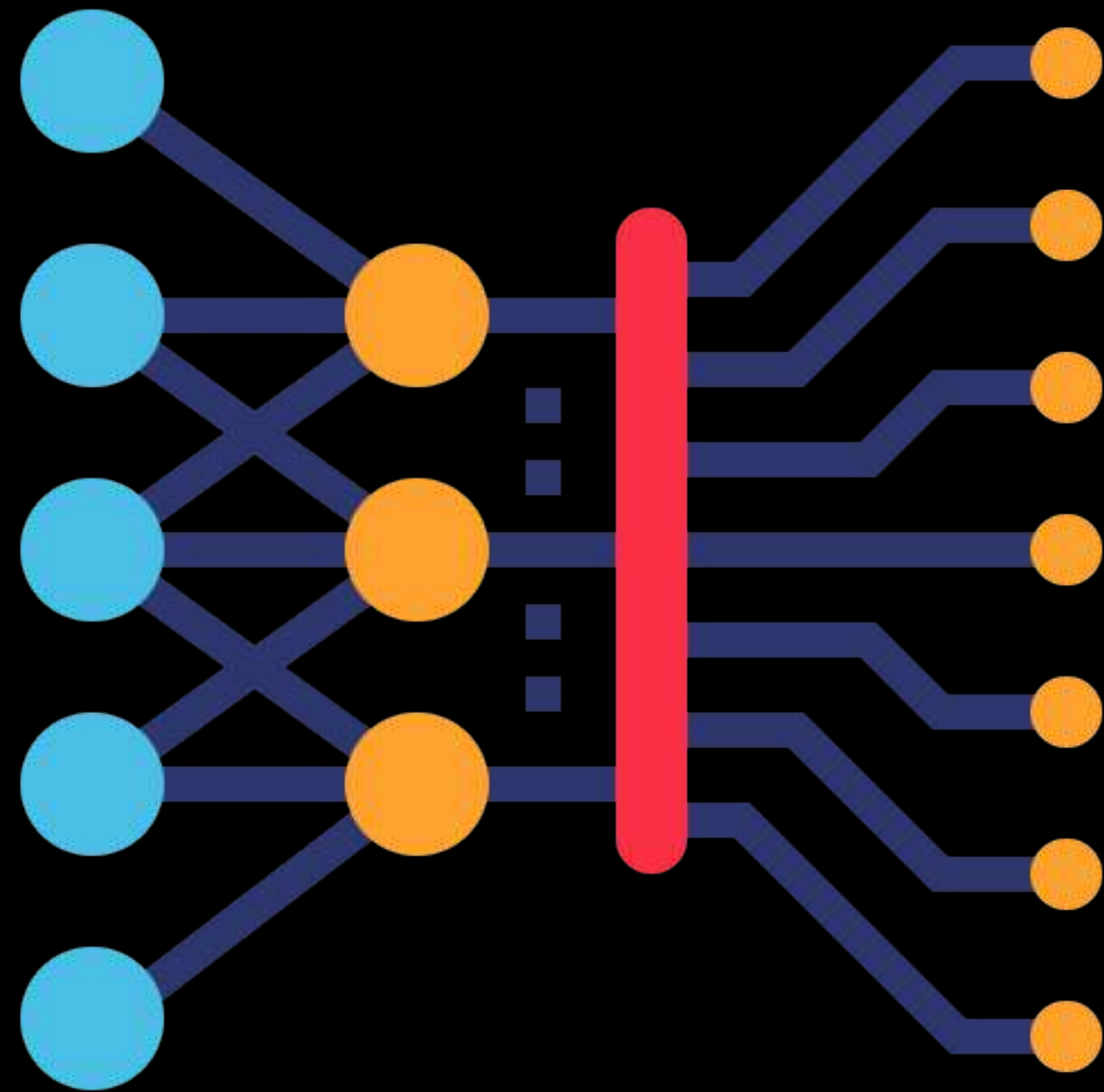
Real-world



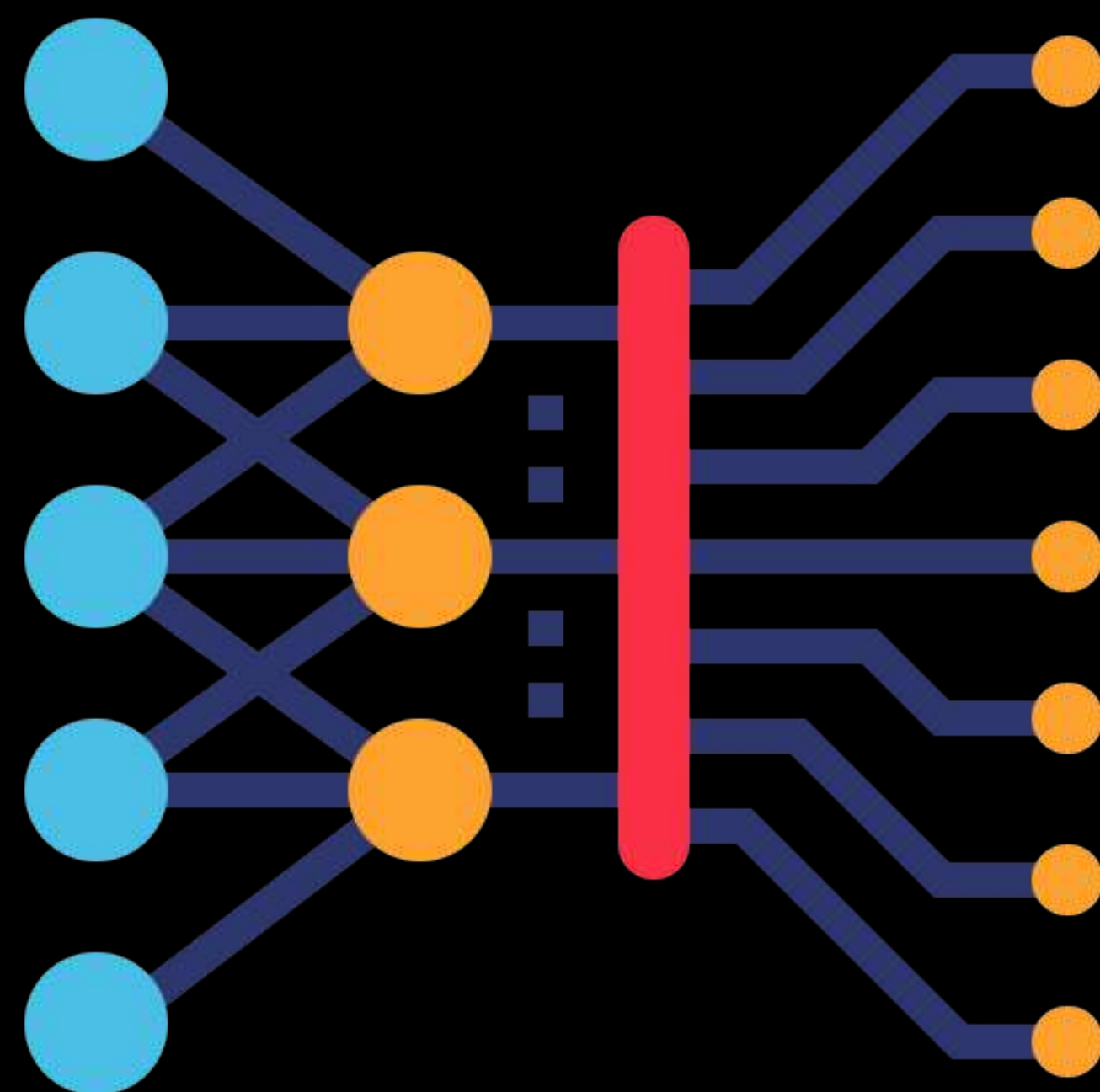
Perception Gap



What is real? How do you define real? If you are talking about what you can feel, you can smell, what you can taste and see, then real is simply electrical signals interpreted by your brain.



What is real? How do you define real? If you are talking about what you can **sense**, you can **perceive**, what you can **handle** and **process**, then real is simply **a feature** interpreted by your **deep neural network**.



sense
perceive
handle process
a feature
deep neural network

Perception Reality Gap

Difference between simulated and real input images

Simulation



Generated

Real-world



Mitigate Gap





Generated



Real-world



Image Quality Metrics

Distribution Level Metrics

IS Inception-score
FID Fréchet Inception Distance
KID Kernel Inception Distance

Single Image Metrics

SSIM Structural Similarity Index
PSNR Peak signal-to-noise ratio
MSE Mean Squared Error
CS Cosine Similarity
TSI Texture Similarity Index
WD Wasserstein Score
KL KL Divergence
HistI Histogram Intersection
CPL Classifier Perceptual Loss

Borji, A. et al. 2018

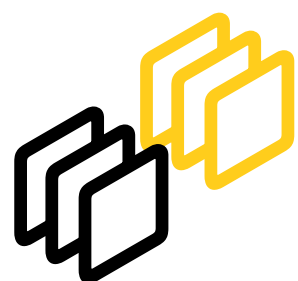
Pang, Y. et al. 2022



Correlation

How do existing Image-to-image evaluation metrics correlate with the associated ADS behaviour?

RQ (Correlation)



Distribution Level Metrics

	Inception-score (IS)		Fréchet Inception Distance (FID)		Kernel Inception Distance (KID)		1	IS and FID are inconsistent across tasks
	Vehicle detection	Lane keeping	Vehicle detection	Lane keeping	Vehicle detection	Lane keeping		
Prediction Error	0.41	0.14	0.24	0.64	0.54	0.54	2	KID is consistent across tasks
Confidence	0.37	0.72	0.21	0.86	0.64	0.74	3	FID is best performer for Lane Keeping, KID for Vehicle detection
Attention Error	0.41	0.65	0.32	0.78	0.86	0.60		
Pearson's correlation coefficient (0,1)								

RQ (Correlation)

Single Image Metrics (2 BEST PERFORMERS)

	Classifier Perceptual Loss (CPL)		Semantic Segmentation Score (SSS)	
	Vehicle detection	Lane keeping	Vehicle detection	Lane keeping
Prediction Error	0.29	0.30	0.23	0.25
Confidence	0.23	0.30	0.21	0.30
Attention Error	×	0.16	×	0.26

- 1

All metrics have weak or negligible correlation
- 2

Multiple metrics have the wrong correlation direction

n

Pearson's correlation coefficient (0,1) [Best of 6 models]

×

At least 1 of 6 datasets has wrong correlation direction



REAL



GENERATED

RQ (Fine-tuning)

Generated



Real-world



Semantic segmentation model

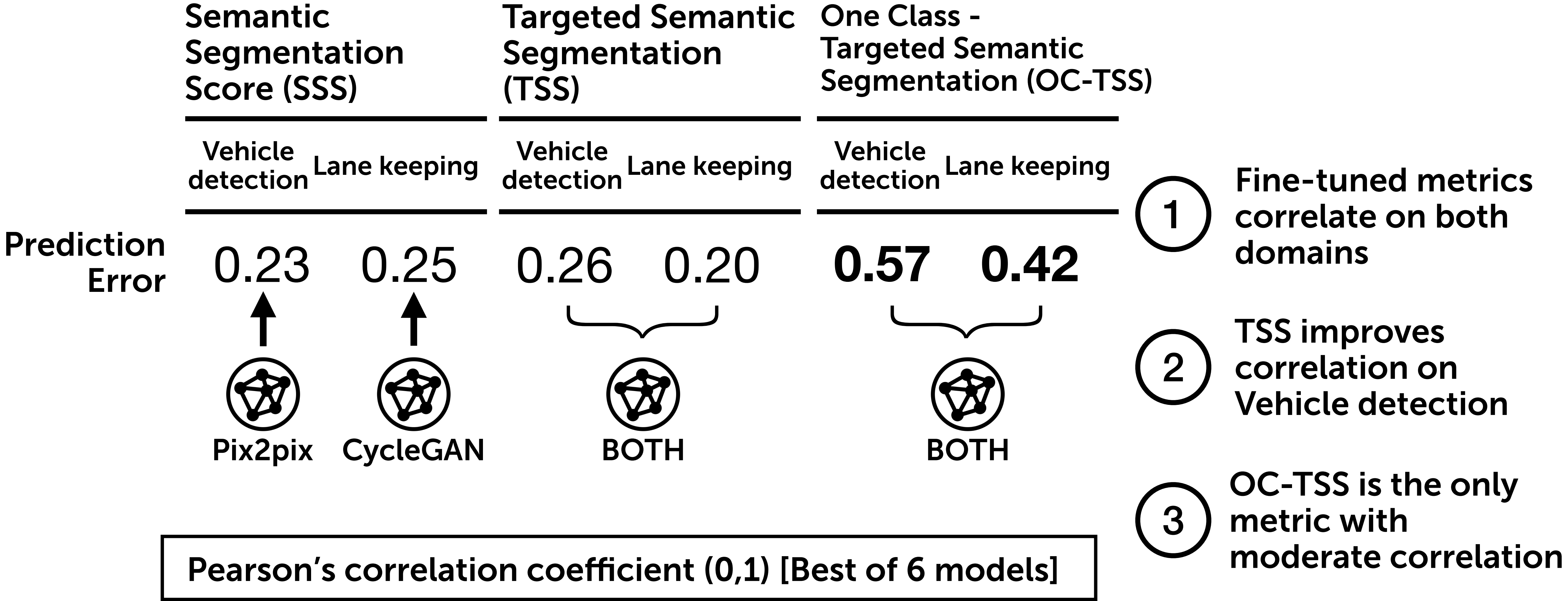
Targeted
Semantic
Segmentation **TSS** =



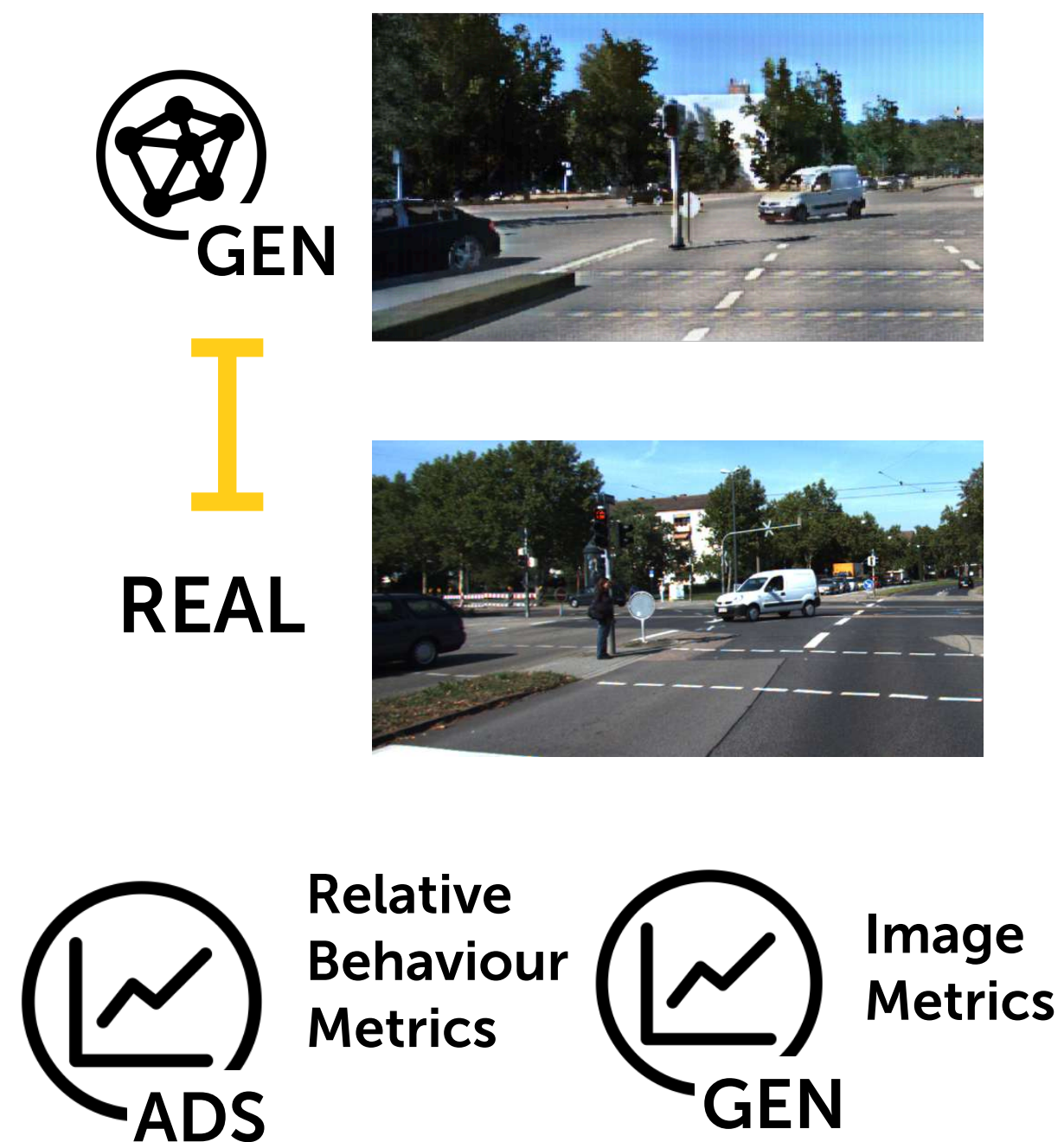
One
Class **OC-TSS** =



RQ (Fine-tuning)



Takeaways

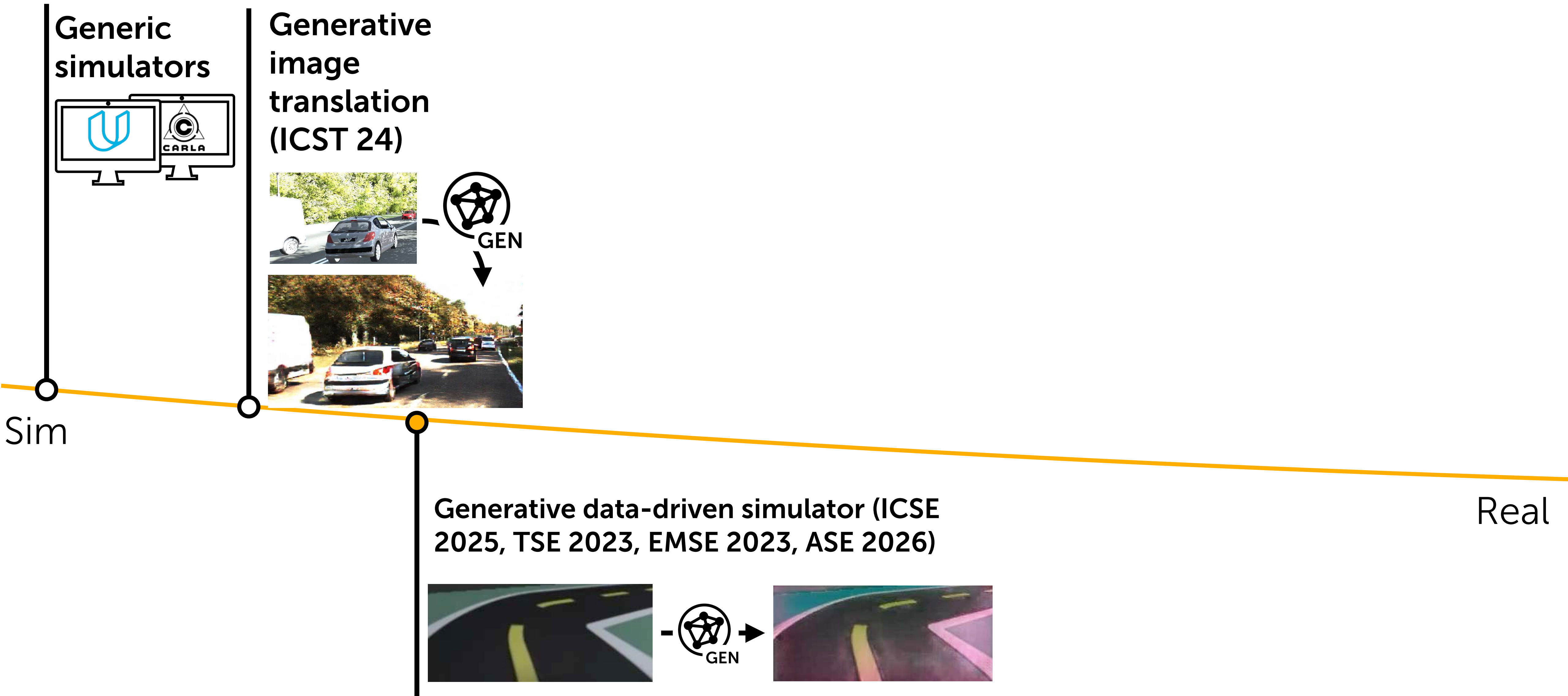


- 1 Image-to-image GenAI tools effectively tackle domain adaptation in ADS
- 2 Current GenAI metrics don't align well with the software behavior that relies on their output
- 3 We need more domain-informed, semantic-aware metrics

Lambertenghi and Stocco.

In Proceedings of 17th IEEE International Conference on Software Testing,
Verification and Validation 2024

Reality Gap Mitigation



ADS requires extensive coverage of the ODD

From regulations to implementation

Existing Standards and Regulations

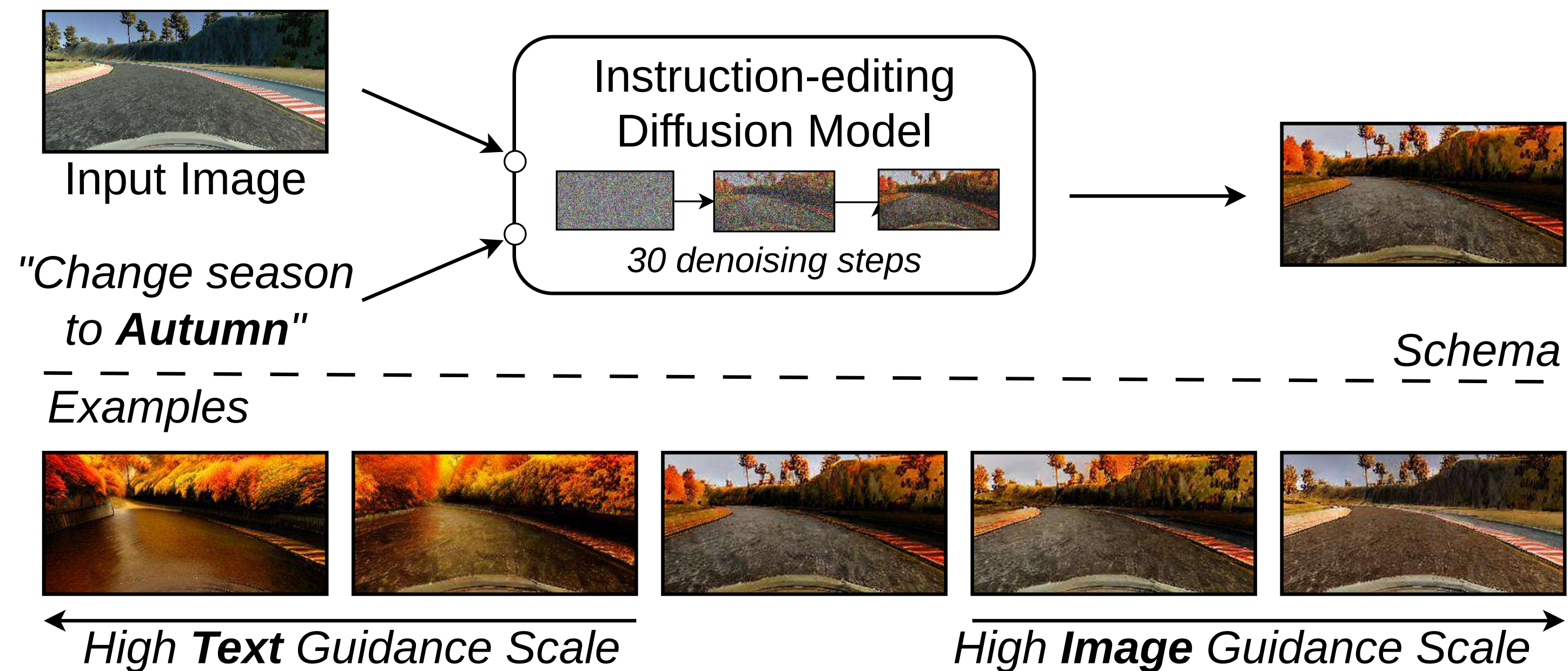
- ISO/PAS 21448 Safety of the Intended Function (SOTIF)
- UN Regulation No 157 (2021/389)
- ISO 34505 “Scenery Elements (Section 9)” and “Environmental Conditions (Section 10)”

Operational Design Domain (ODD)

- roadway types
- geographic area
- speed range
- environmental conditions (weather as well as day/night time)

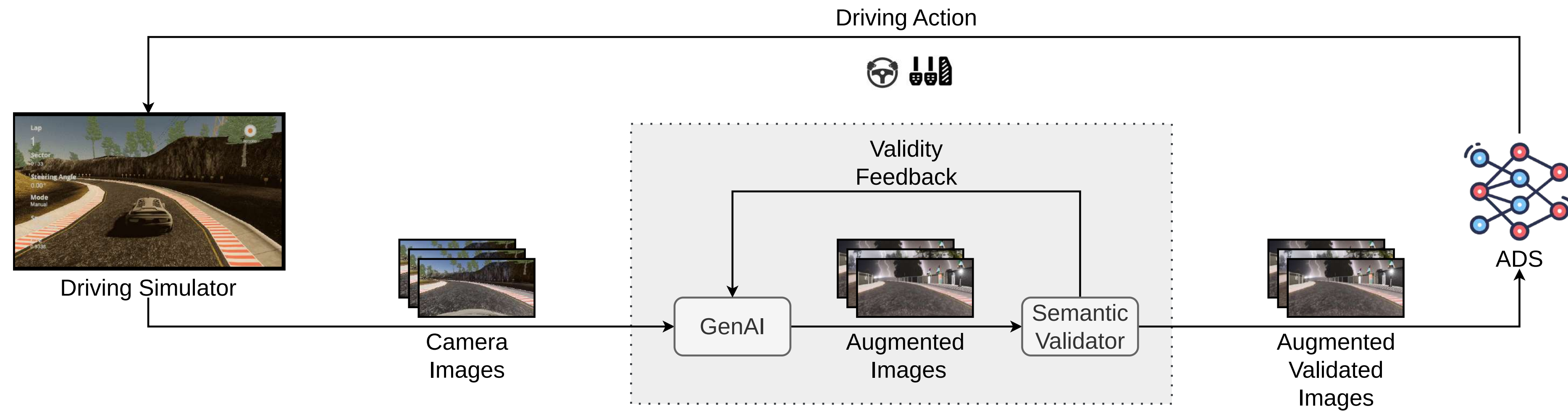
Instruction-editing Diffusion Models

Prompt: Textual



Enhancing ADS Testing with Driving Simulators and Generative AI

Simulators with Generative AI



Enhancing ADS Testing with Driving Simulators and Generative AI

Simulators with Generative AI (naïve integration)

InstructPix2Pix



- ✓ Diversity
- ✗ Temporal Consistency
- ✗ Efficiency

Enhancing ADS Testing with Driving Simulators and Generative AI

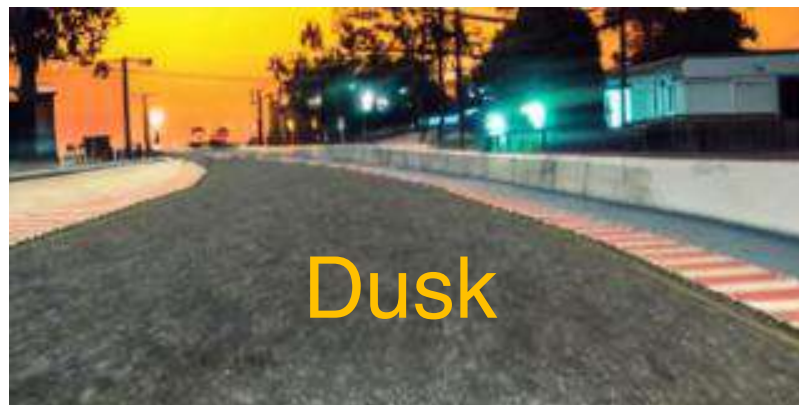
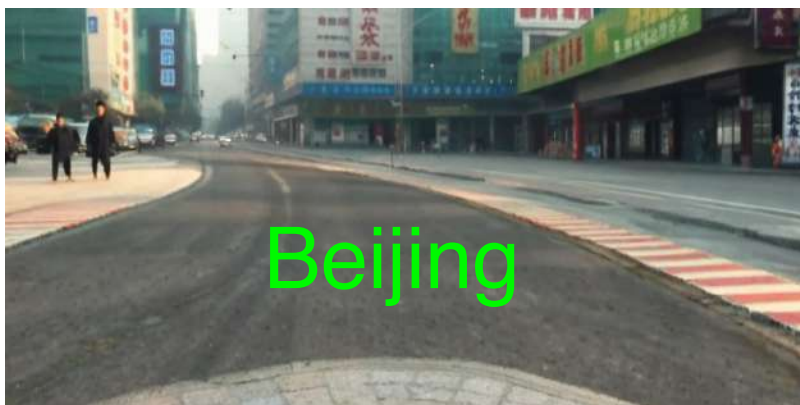
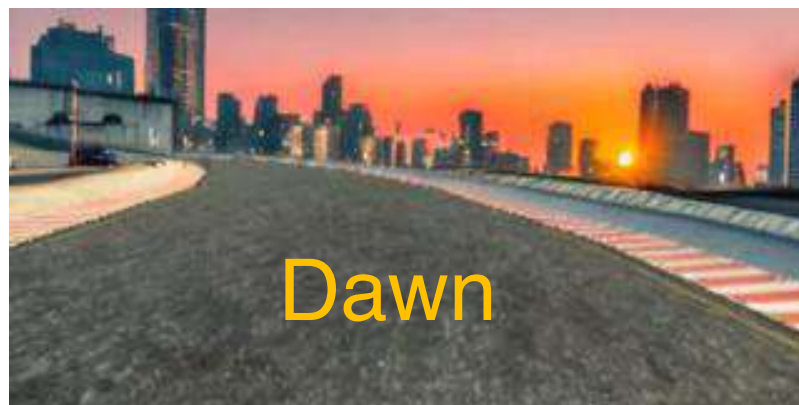
Simulators with Generative AI (knowledge distillation)

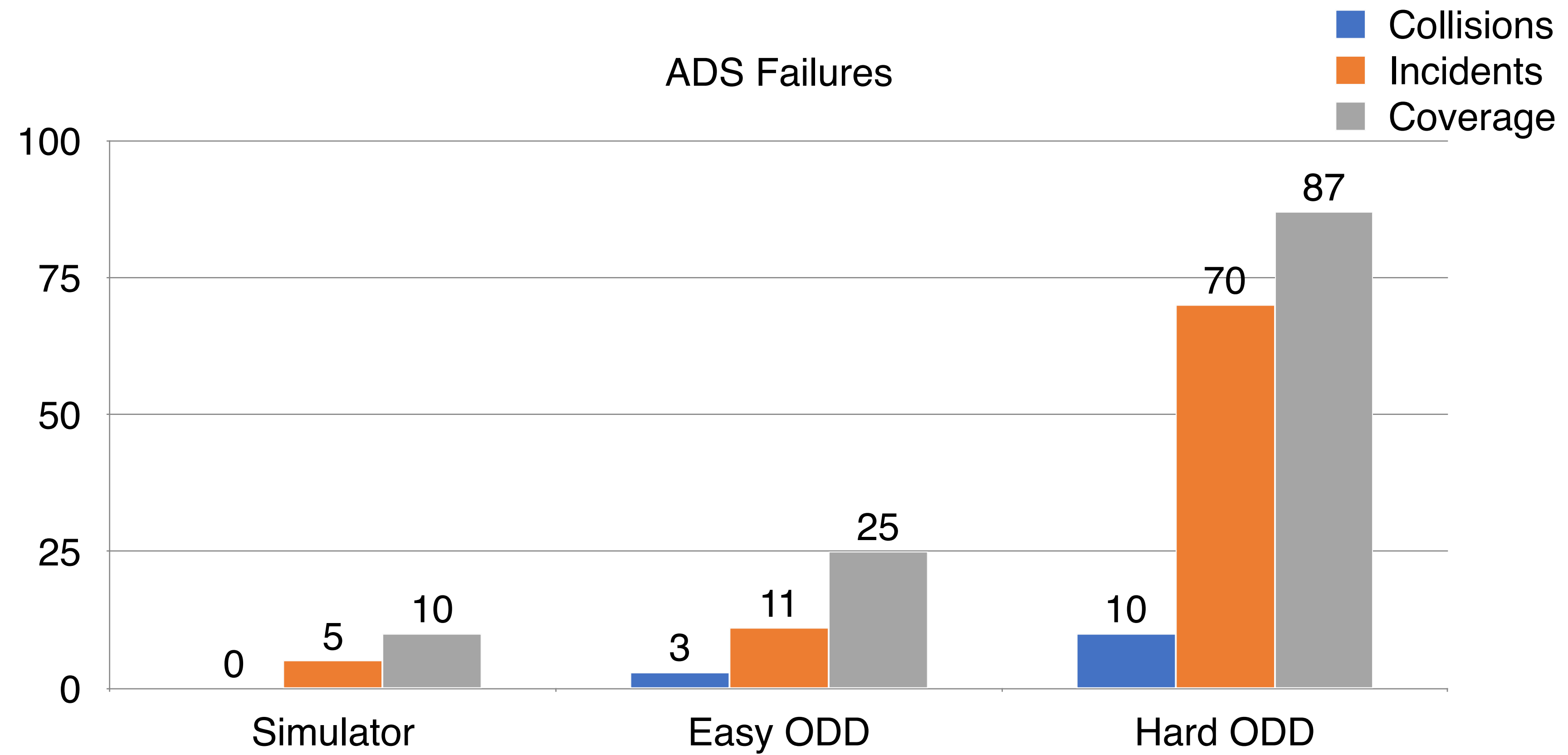
InstructPix2Pix with Knowledge Distillation



- ✓ Diversity
- ✓ Temporal Consistency
- ✓ Efficiency







Takeaways

ORIG
ODD



NEW
ODD



Behaviour
Metrics



Diffusion
Models

①

Diffusion models effectively tackle domain generation for ADS testing

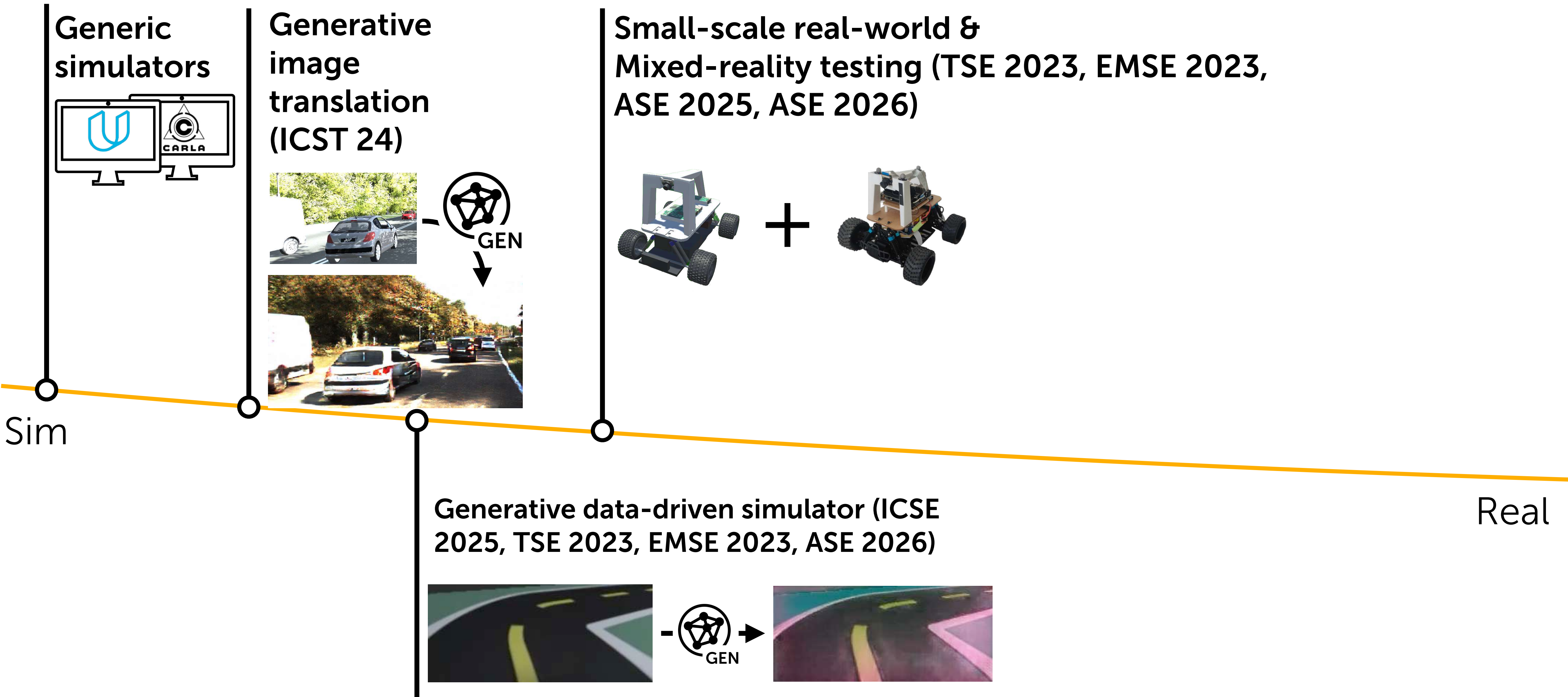
②

They complement simulator testing, uncovering failures in areas previously considered error-free

③

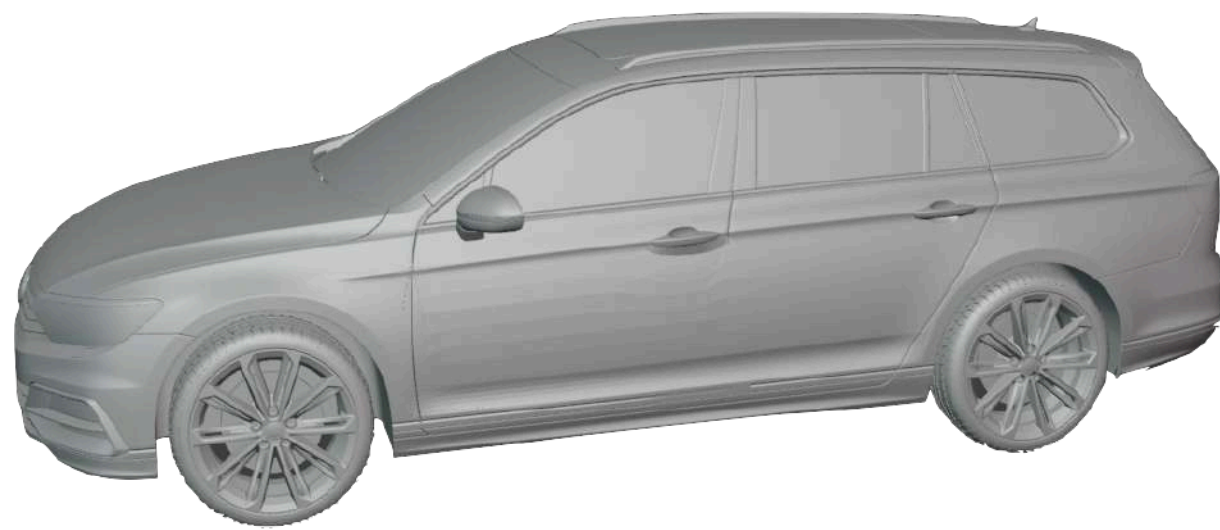
Knowledge distillation is key to achieving high simulation efficiency

Reality Gap Mitigation



Gap mitigation

Simulation testing



Simulated

Behaviour
Scenarios
Sensors

Small-scale Real-world testing



Real

Behaviour
Sensors

Full-scale Real-world testing

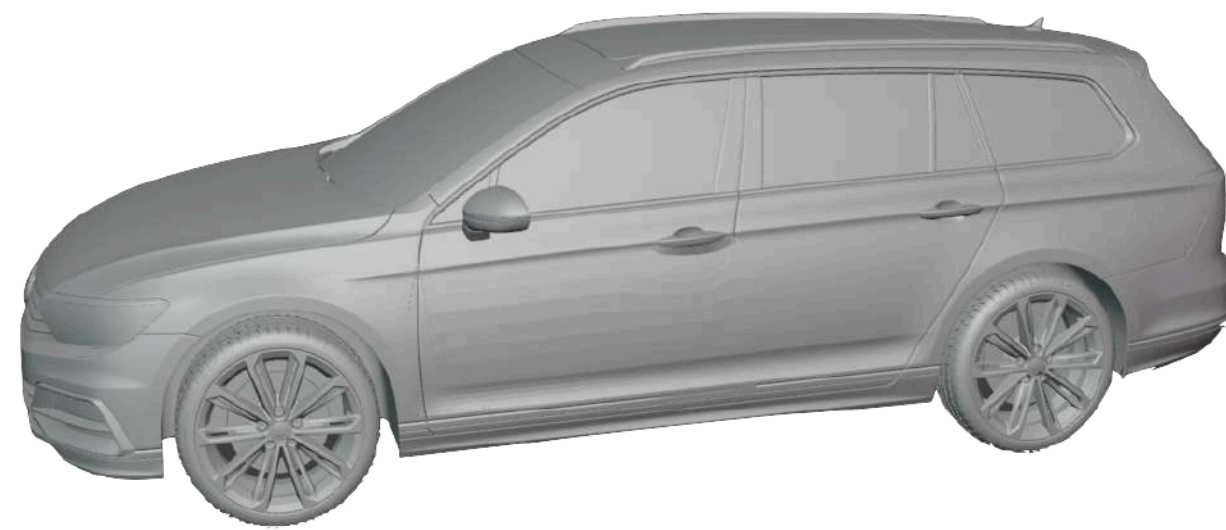


Real

Behaviour
Scenarios
Sensors

Gap mitigation

Simulation testing



Simulated

Behaviour
Scenarios
Sensors

Small-scale Real-world testing



Real

Behaviour
Sensors

Full-scale Real-world testing



Real

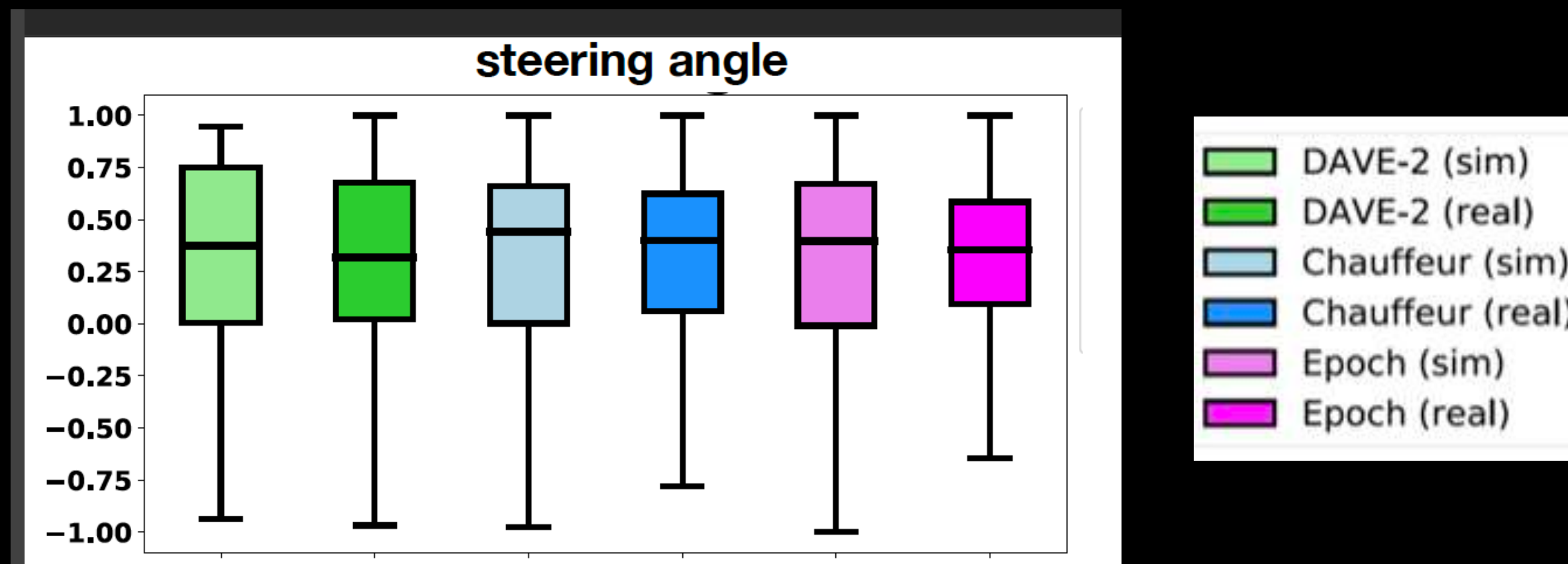
Behaviour
Scenarios
Sensors

Same lane-keeping model

architecture, two worlds

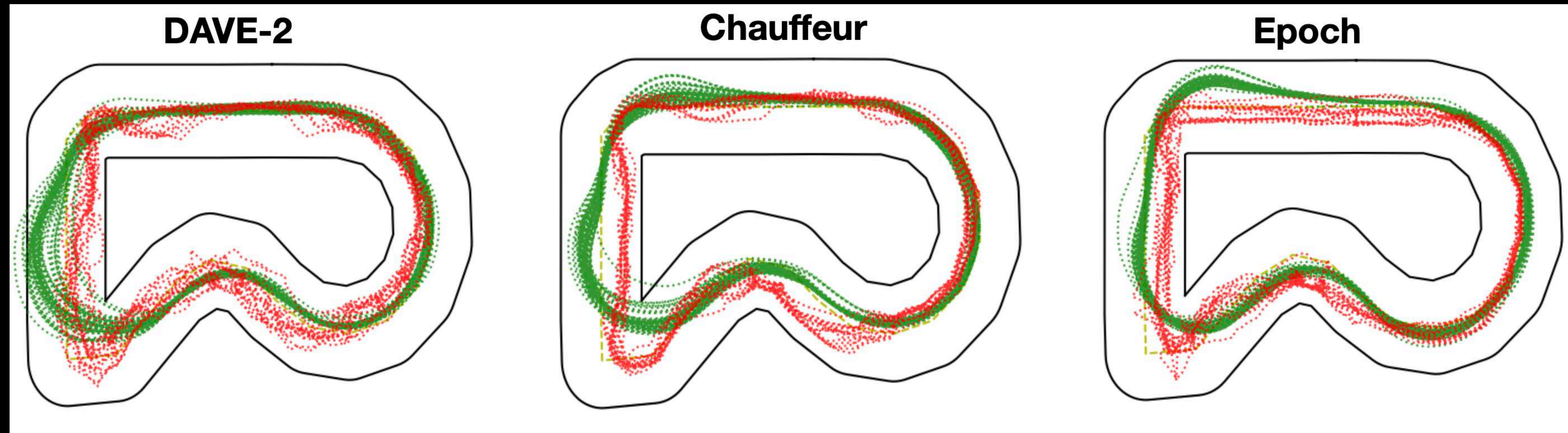


Would the same driving model behave the same?



Steering angle distributions **do transfer** across simulated and real-world environments

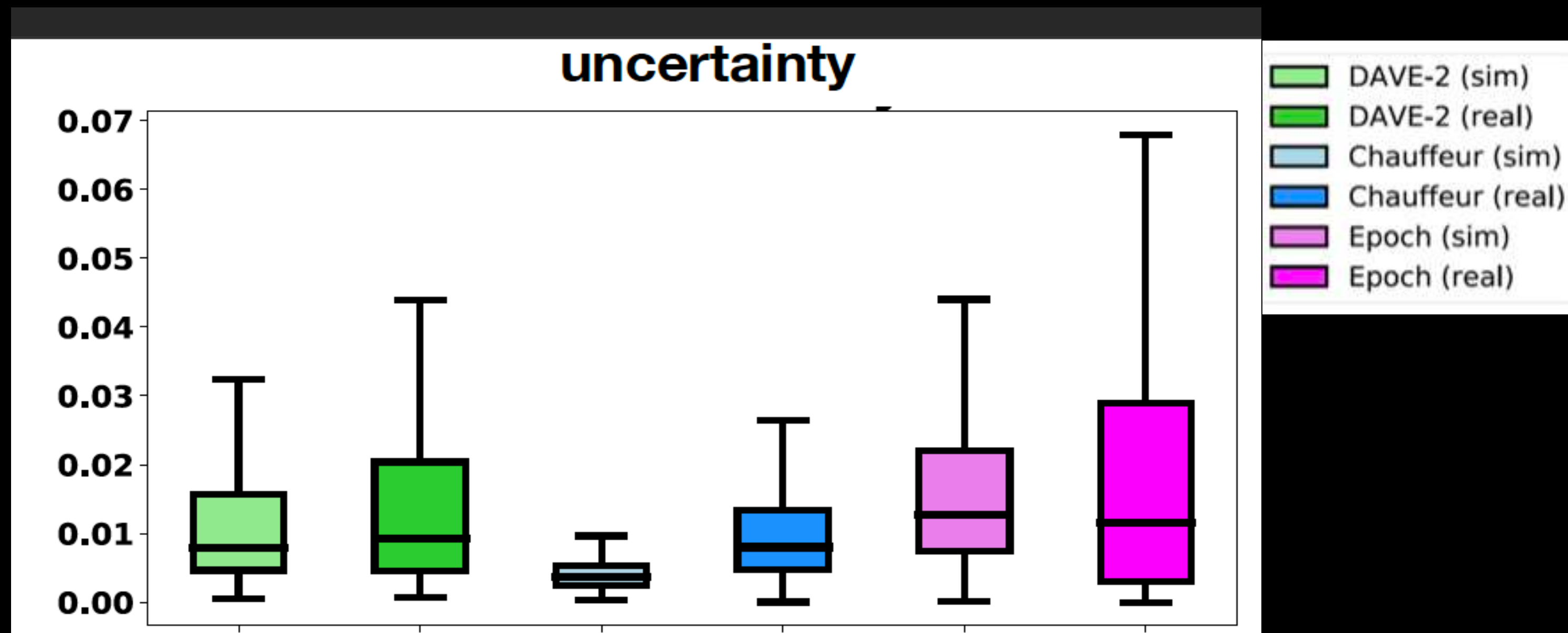
Expected as we aligned the two environments.
It may suggest that component-level testing
is an option but...



Virtual (green) and physical (red) trajectories.

Lateral position is different across simulated and real-world environments

Component-level testing is not an option,
we need system-level testing

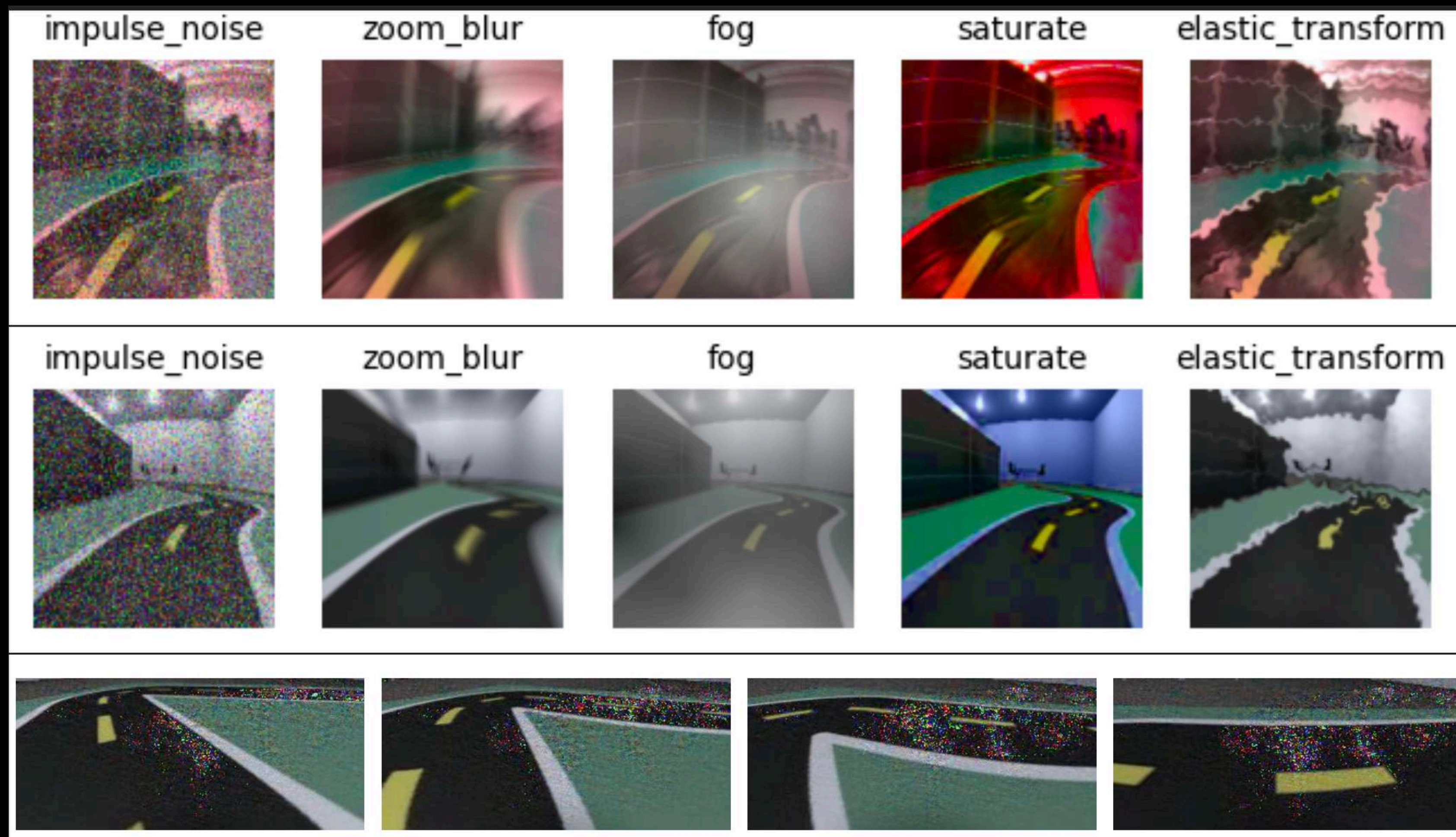


Uncertainty is higher in real-world environments

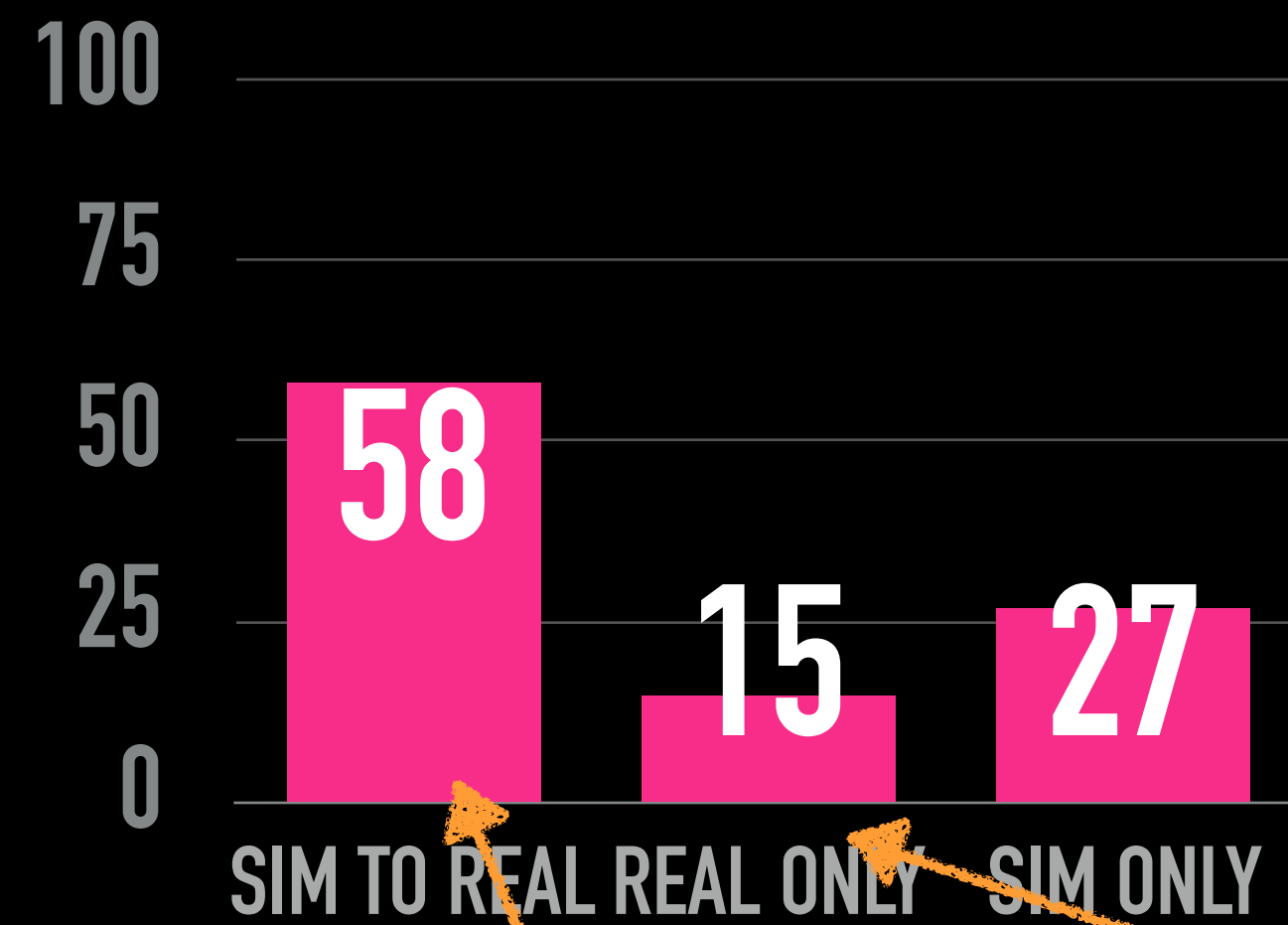
We need real-world testing (or better simulators)!

Would it fail the same?

We test both simulated and real cars under the same conditions (corruptions and adv examples)

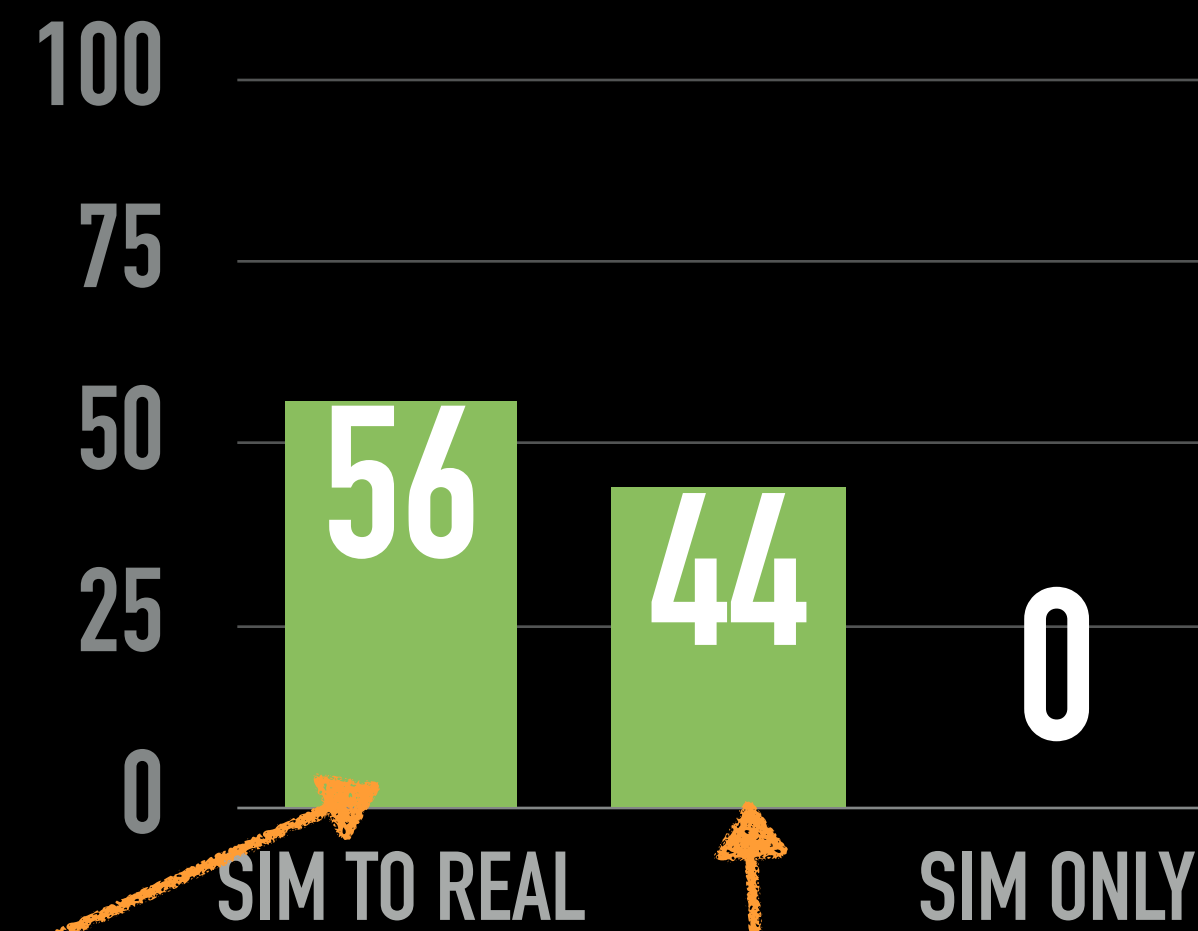


CORRUPTIONS



Most testing
results
transfer

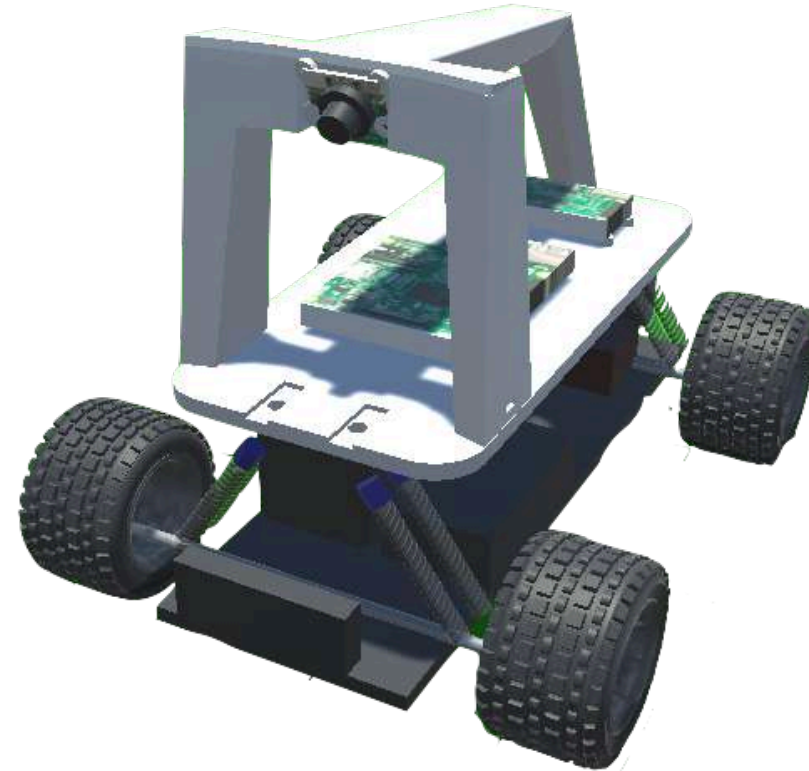
ADV. EXAMPLES



We cannot
completely avoid
real-world testing

Gap mitigation

**Simulation
testing**



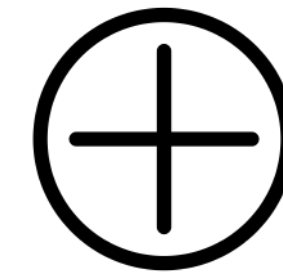
+

**Real-world
testing**

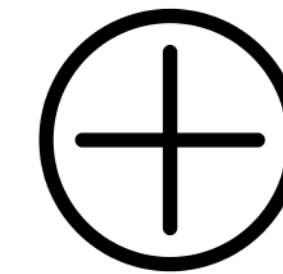


=

Mixed-reality



More realistic than sim



More versatile than full-scale

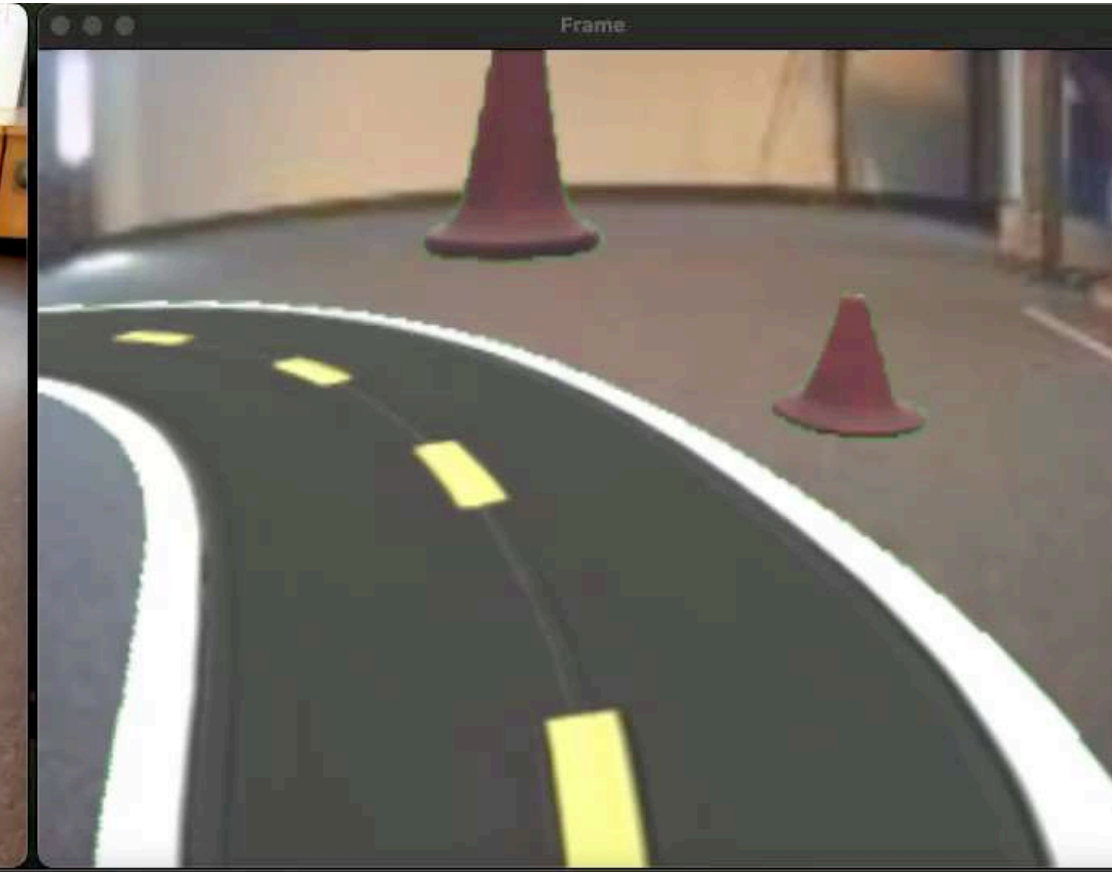
Mixed-reality framework

○ End-to-end ADS execution

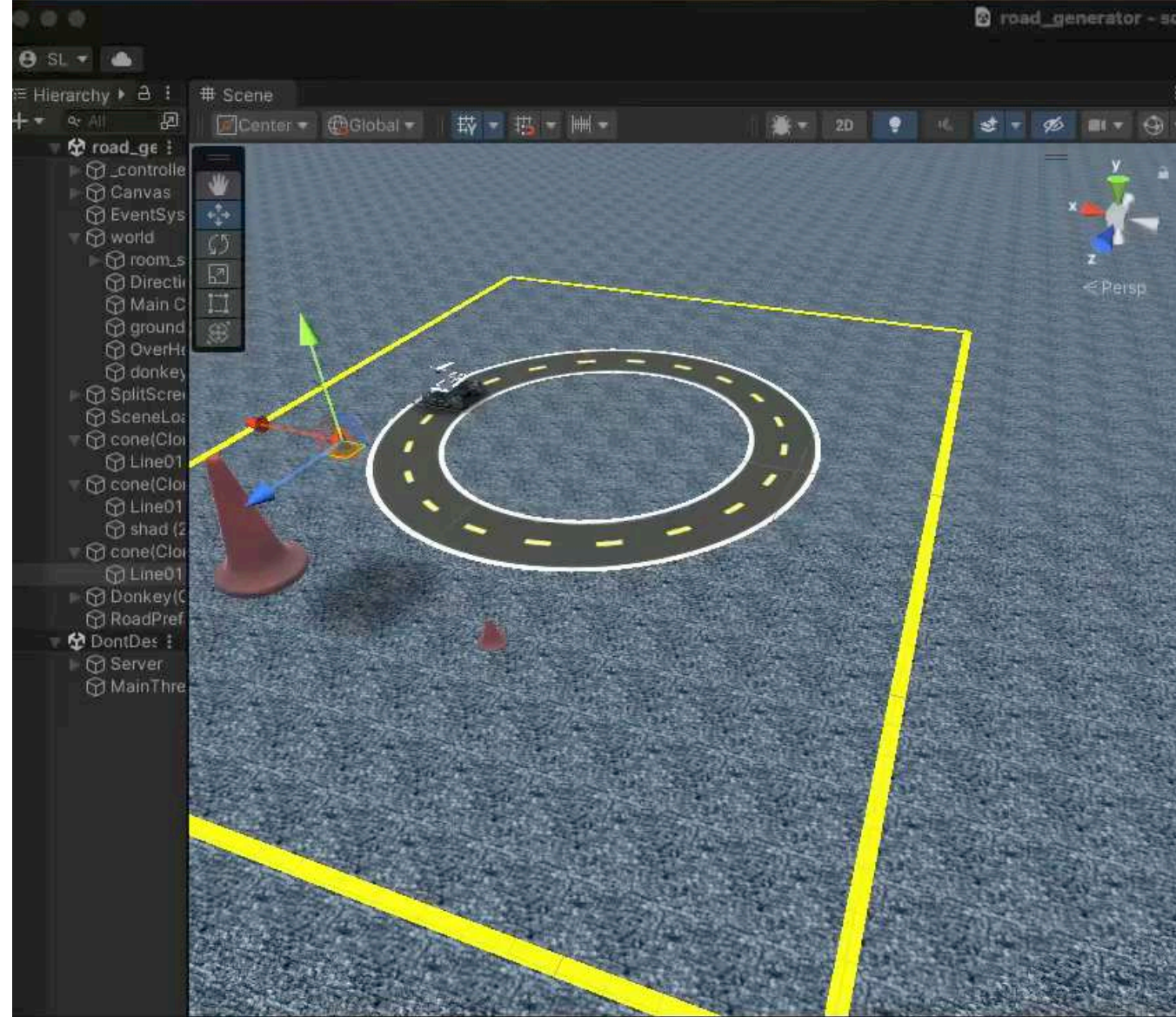
Real world



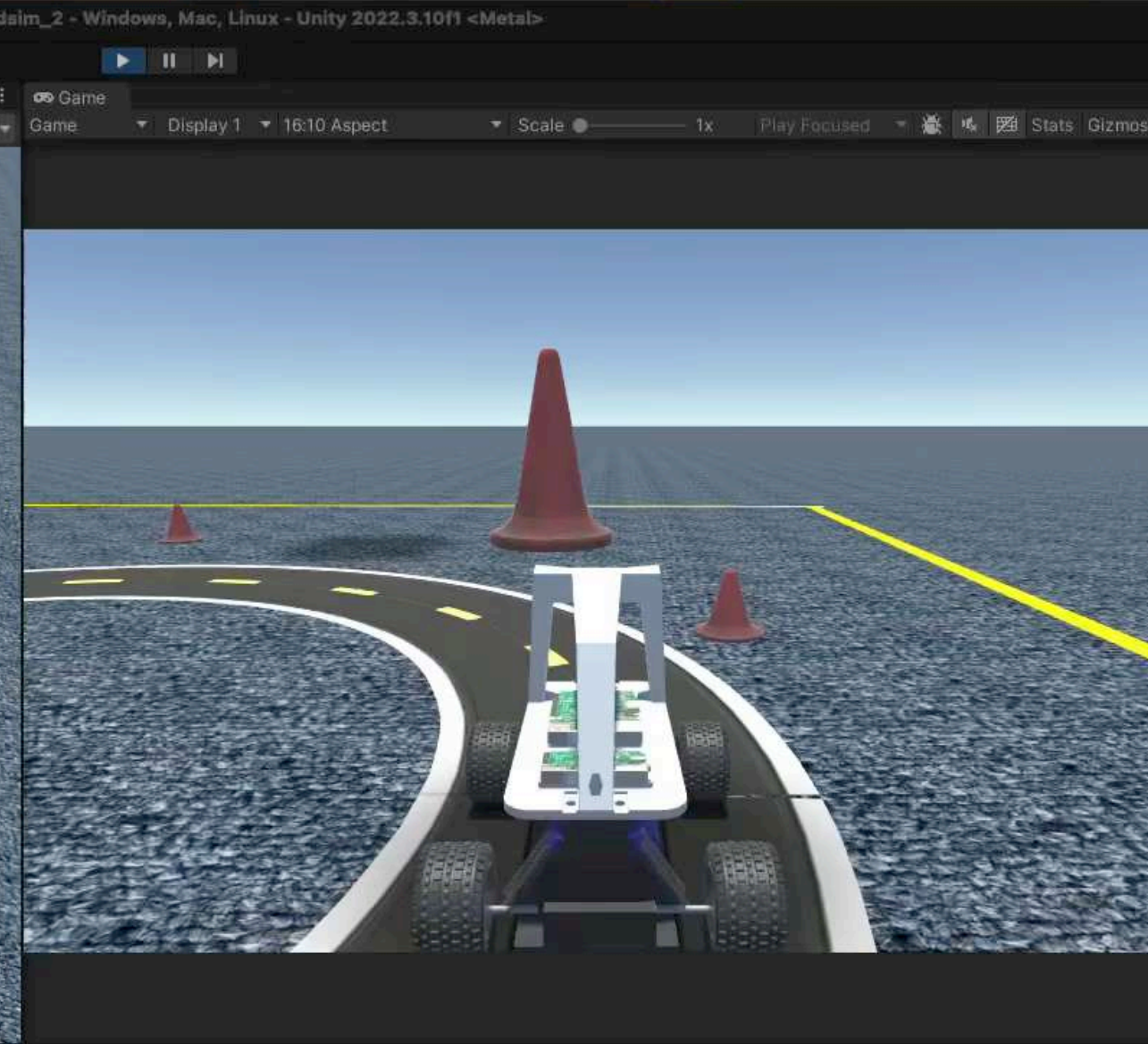
Mixed-reality



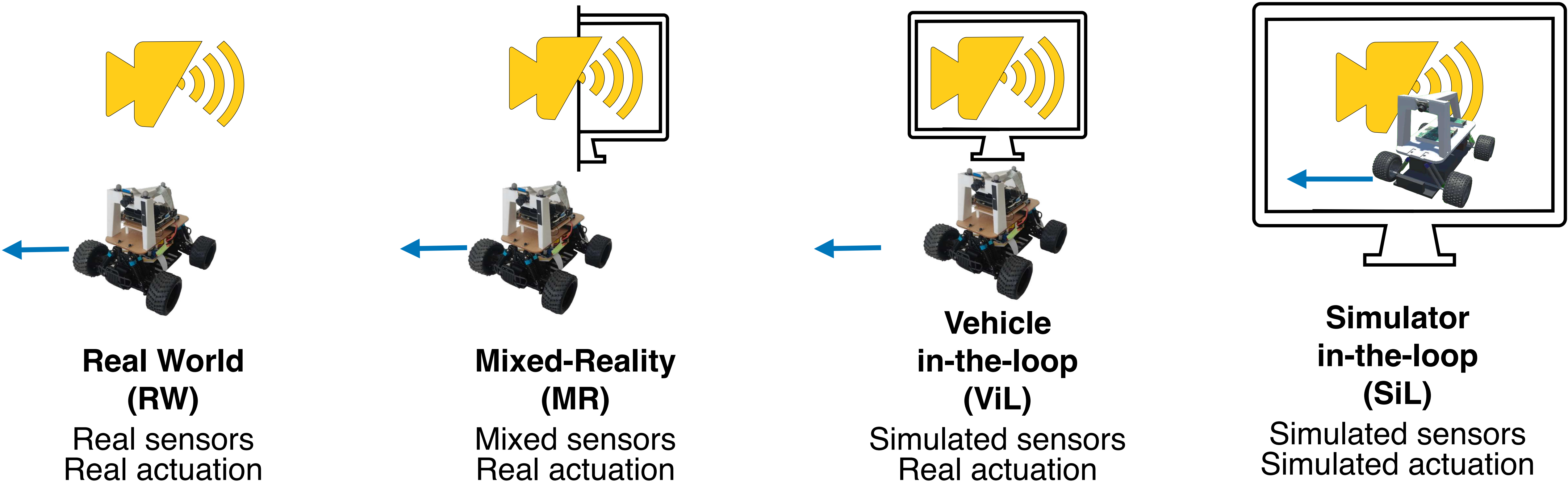
Simulation



Simulation
(FOV)



Testing modalities



Our
GOAL

Understand the ability of SiL, ViL, and MR to mitigate the actuation gap & perception gap

Testing modalities

RW



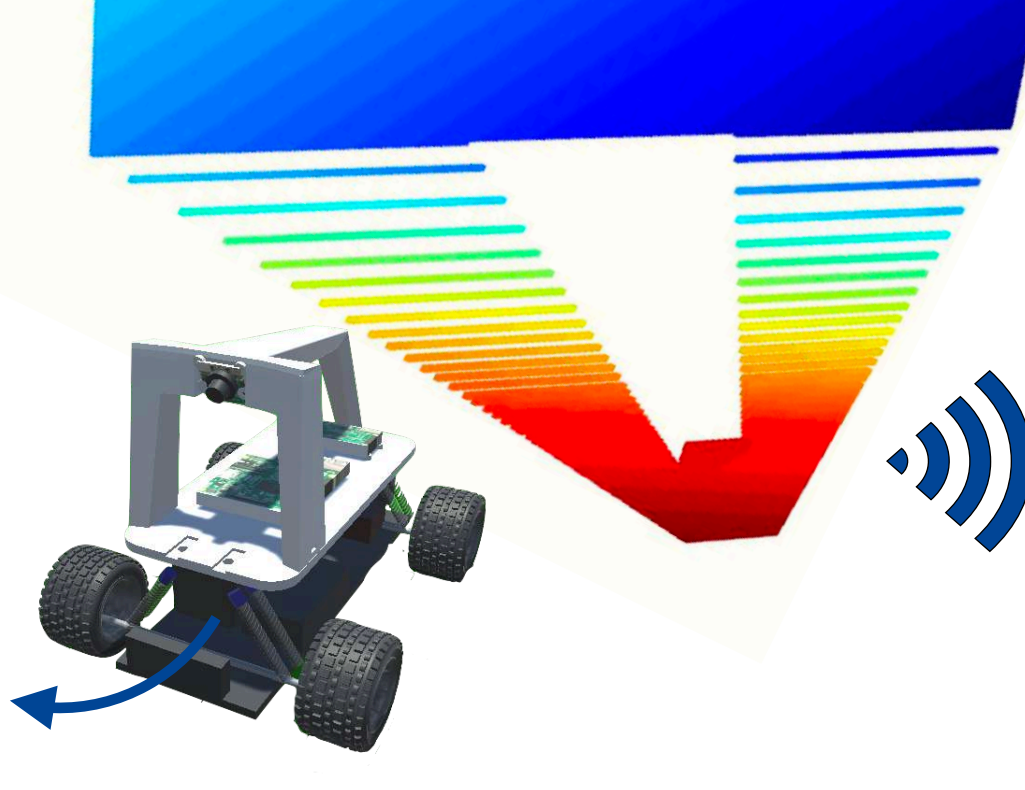
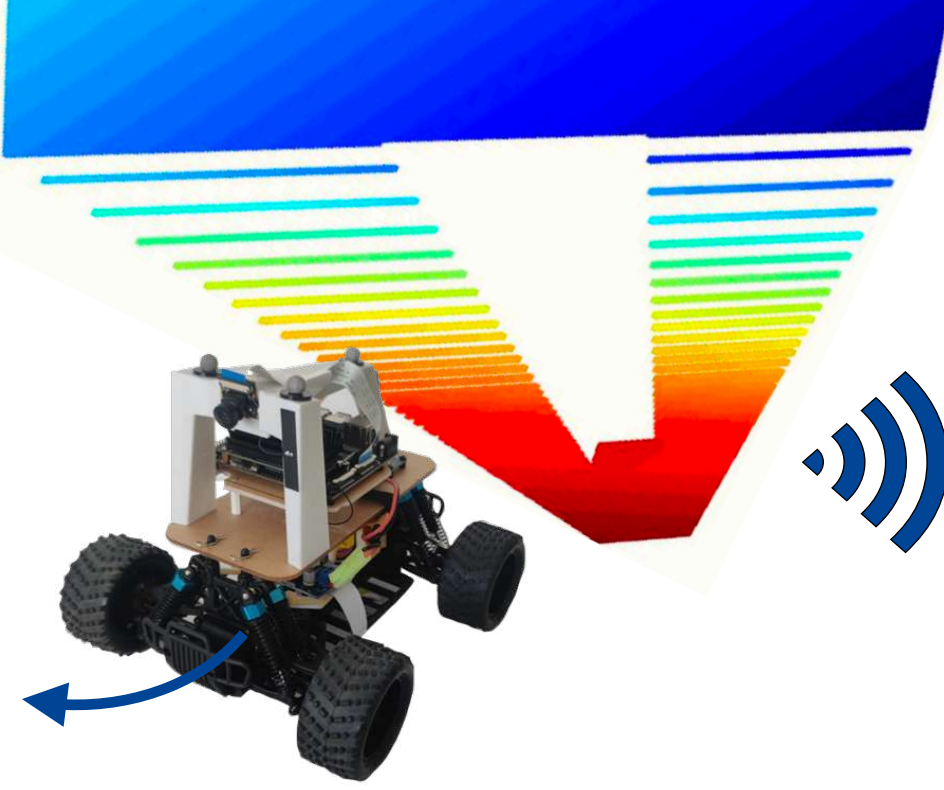
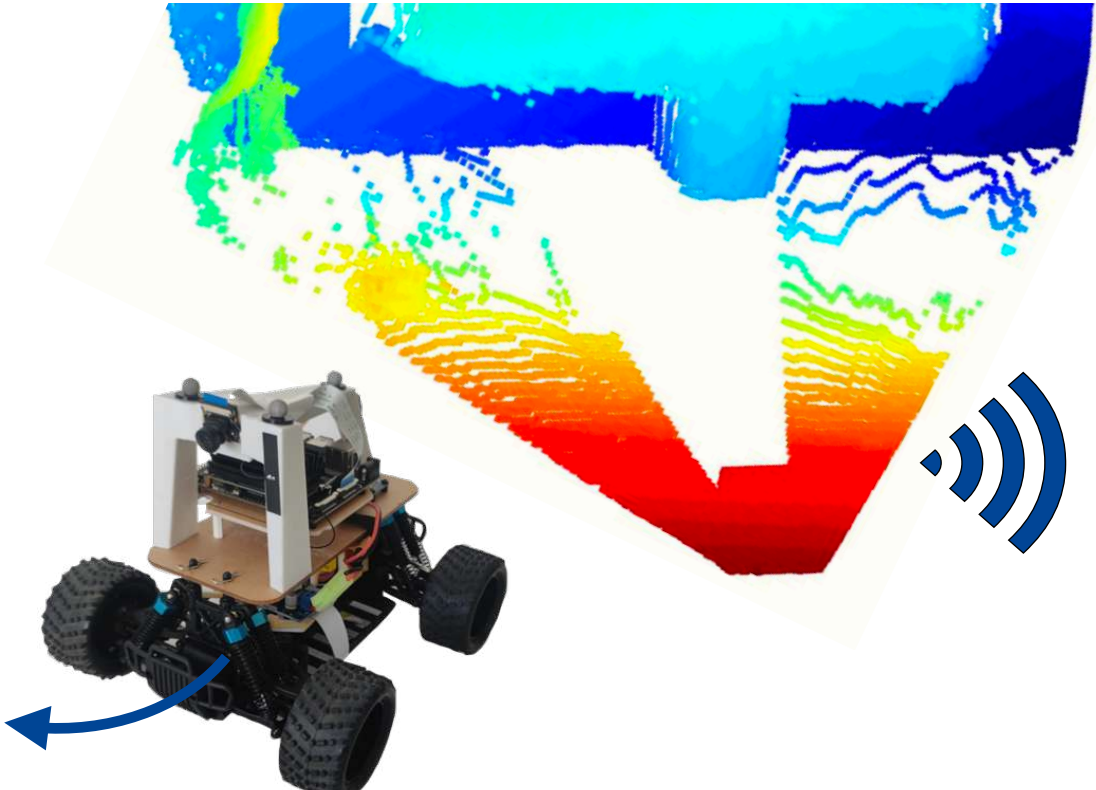
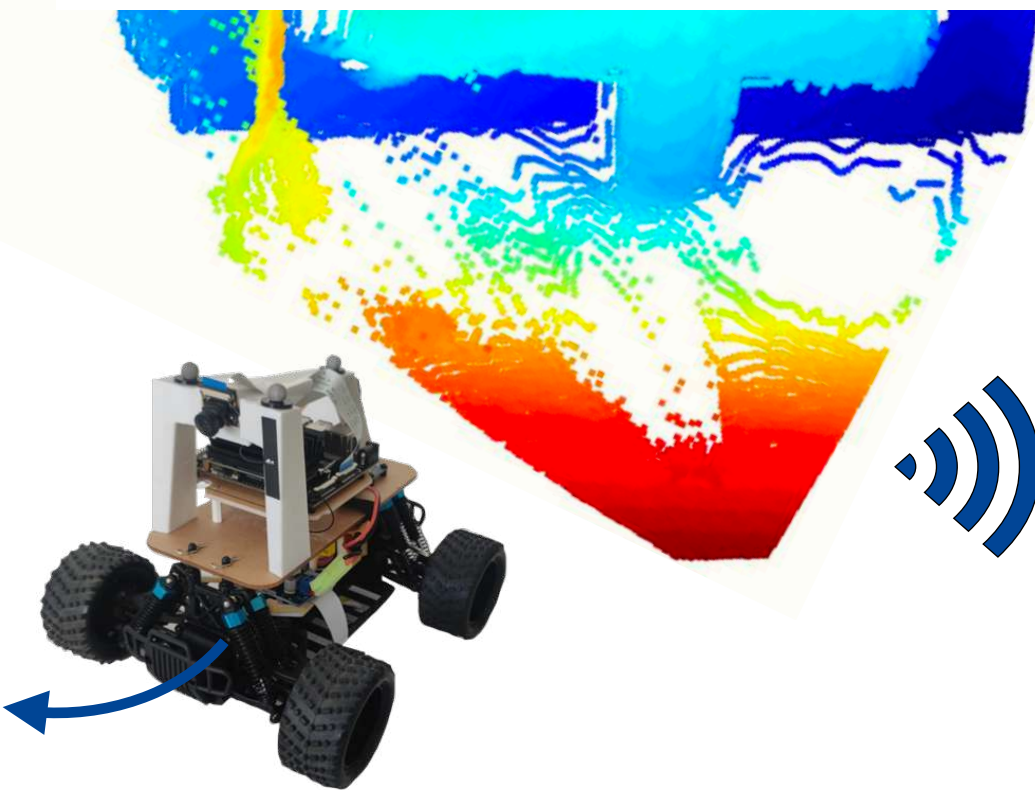
MR



ViL



SiL



**Real World
(RW)**

Real sensors
Real actuation

**Mixed-Reality
(MR)**

Mixed sensors
Real actuation

**Vehicle
in-the-loop
(ViL)**

Simulated sensors
Real actuation

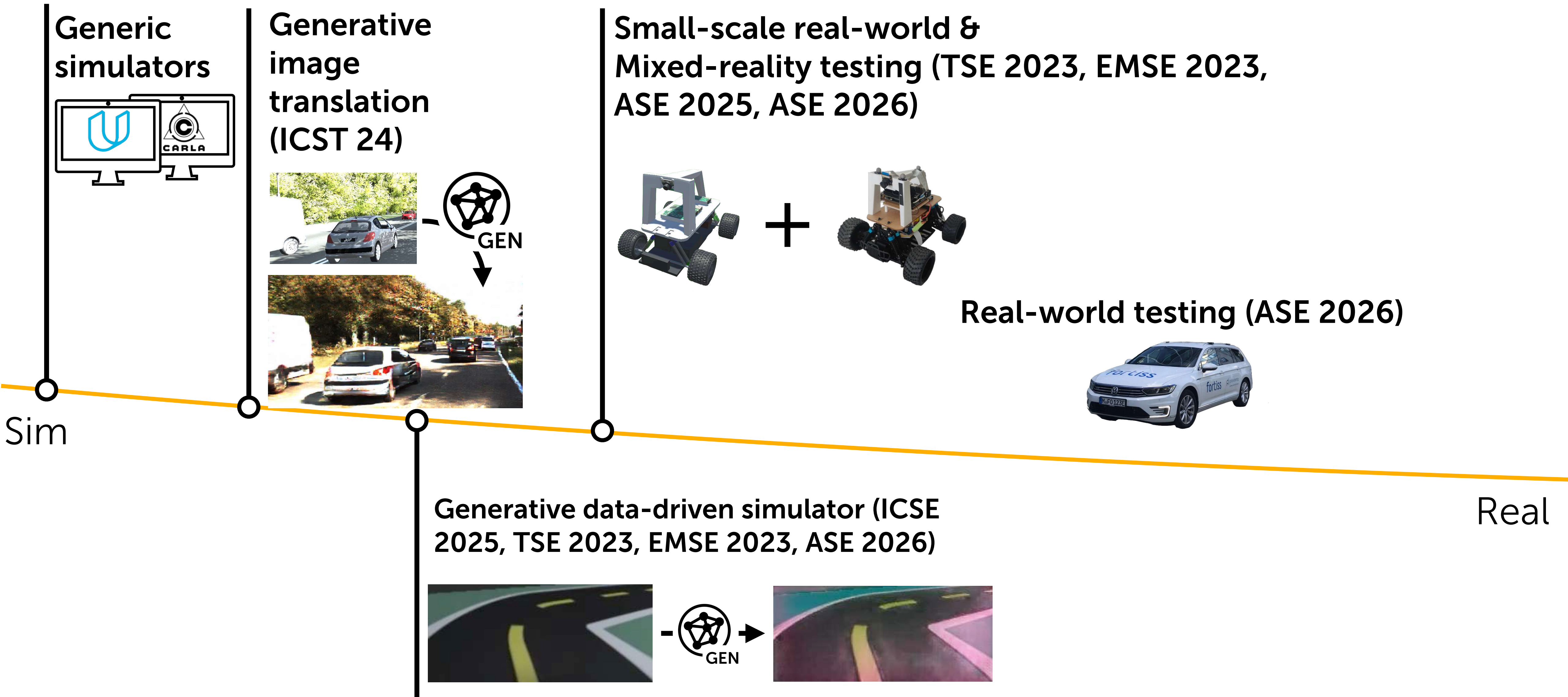
**Simulator
in-the-loop
(SiL)**

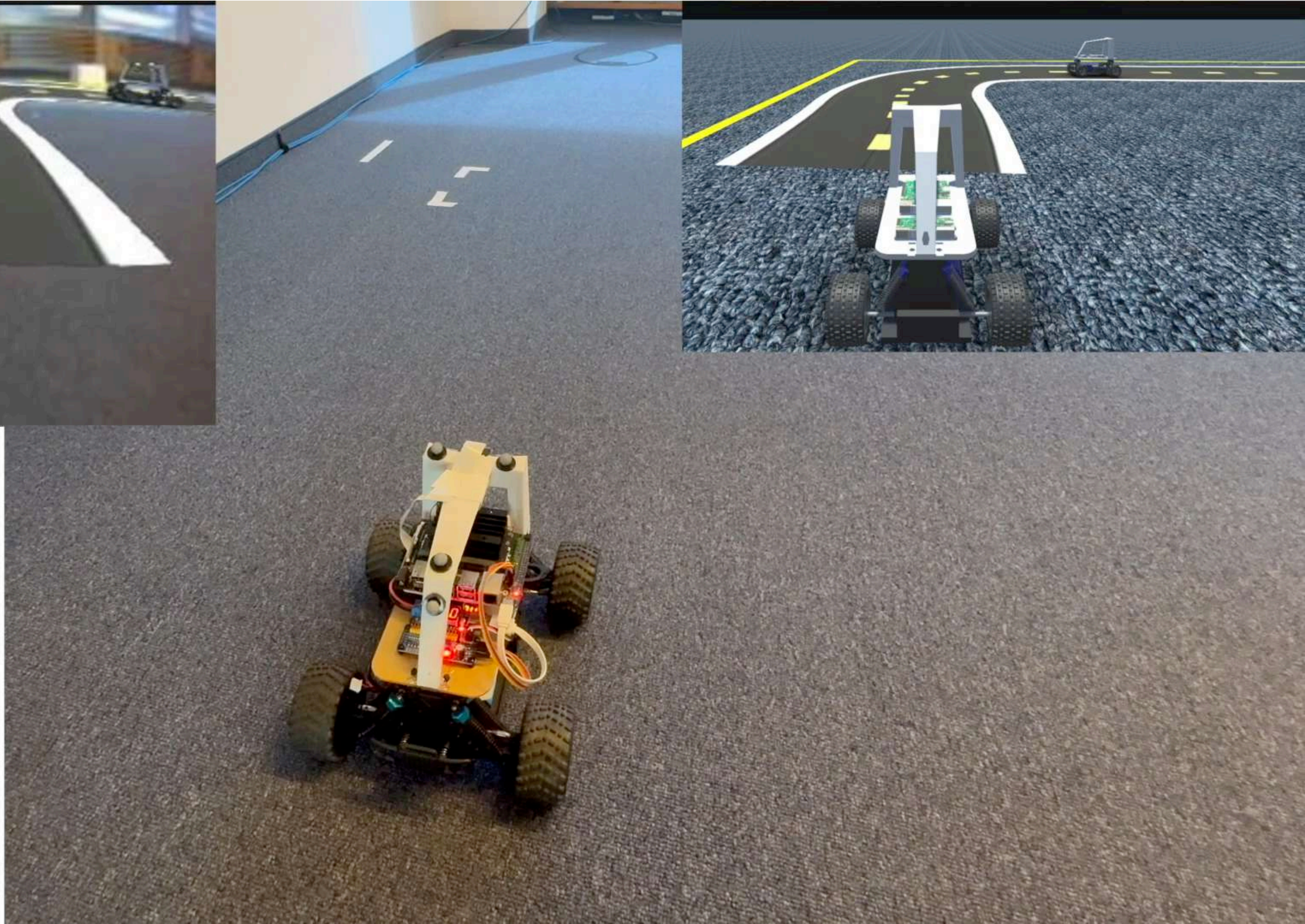
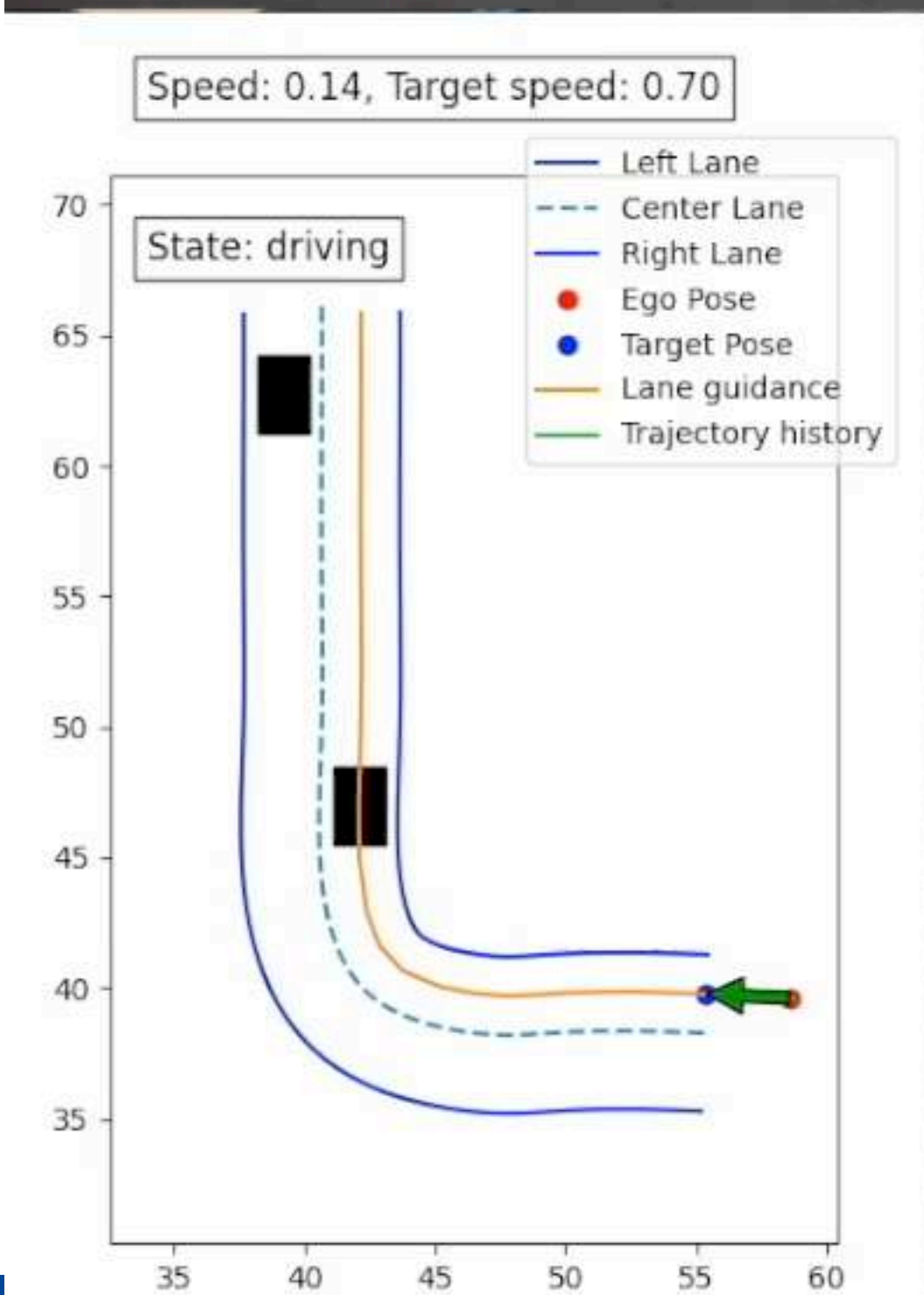
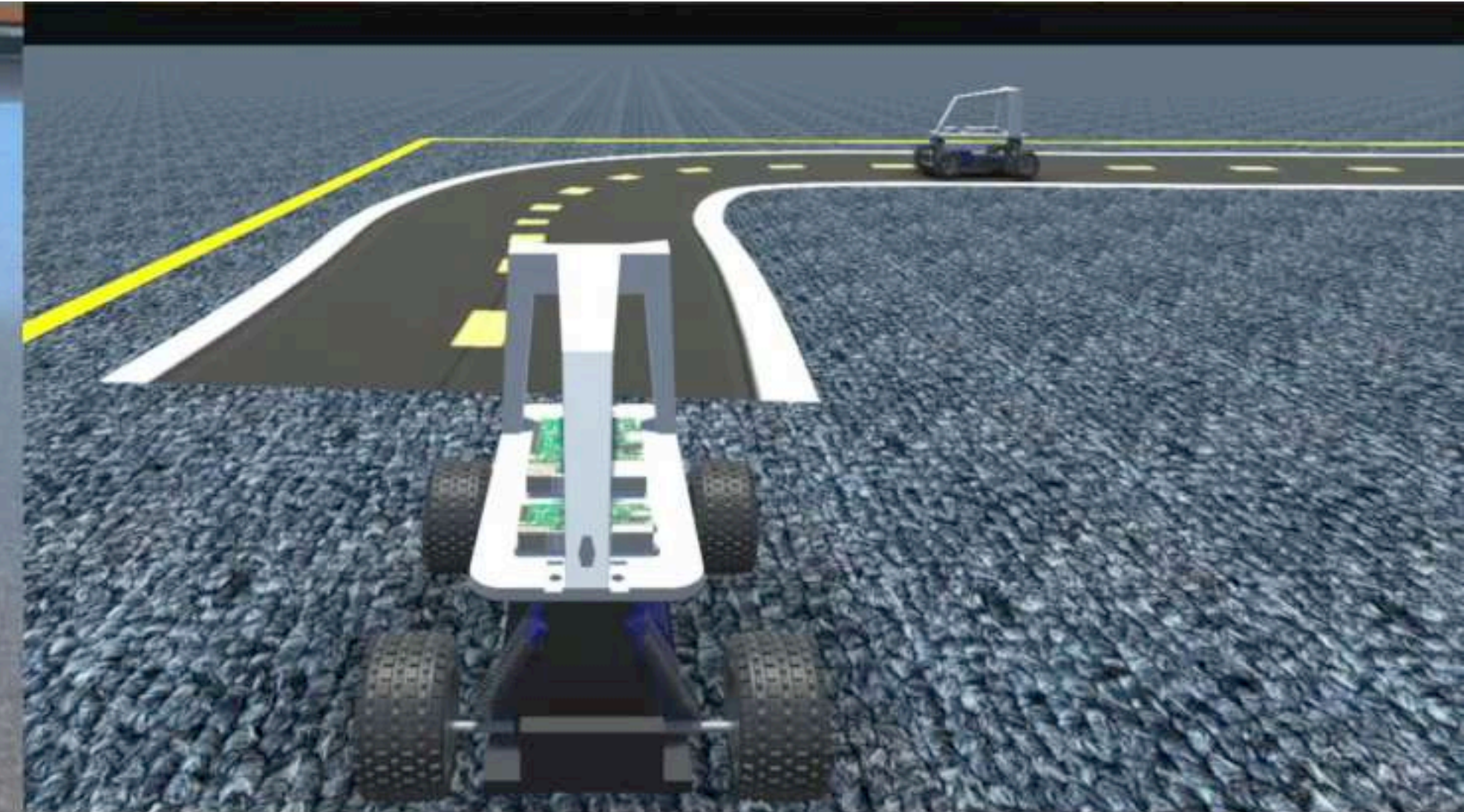
Simulated sensors
Simulated actuation

Takeaways

- ① **MR best matches real behaviour, ViL closes the actuation gap to real dynamics, SiL misrepresents outcomes due to oversimplified dynamics**
- ② **Evaluation is architecture-dependent and unreliable**
 - ▶ **E2E systems appear worse in simulation (under-confident, fail in SiL but succeed in RW)**
 - ▶ **Modular systems appear better in simulation (over-confident, succeed in SiL but fail in RW), hence the same simulator gives opposite conclusions**
- ③ **The perception gap, not actuation, is the dominant source of failure**

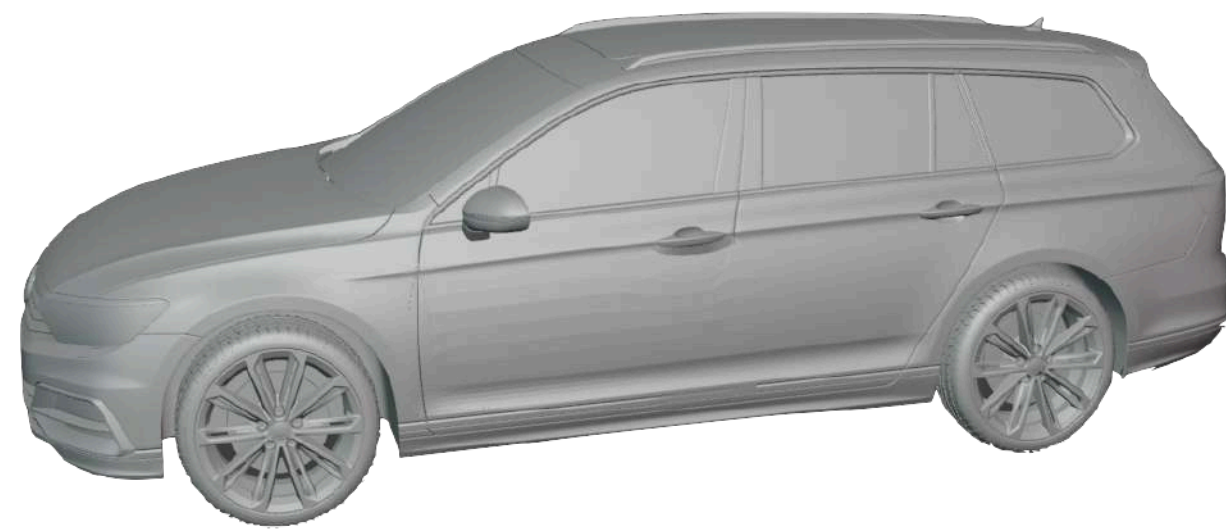
Reality Gap Mitigation





Gap mitigation

Simulation testing



Simulated

Behaviour
Scenarios
Sensors

Small-scale Real-world testing



Real

Behaviour
Sensors

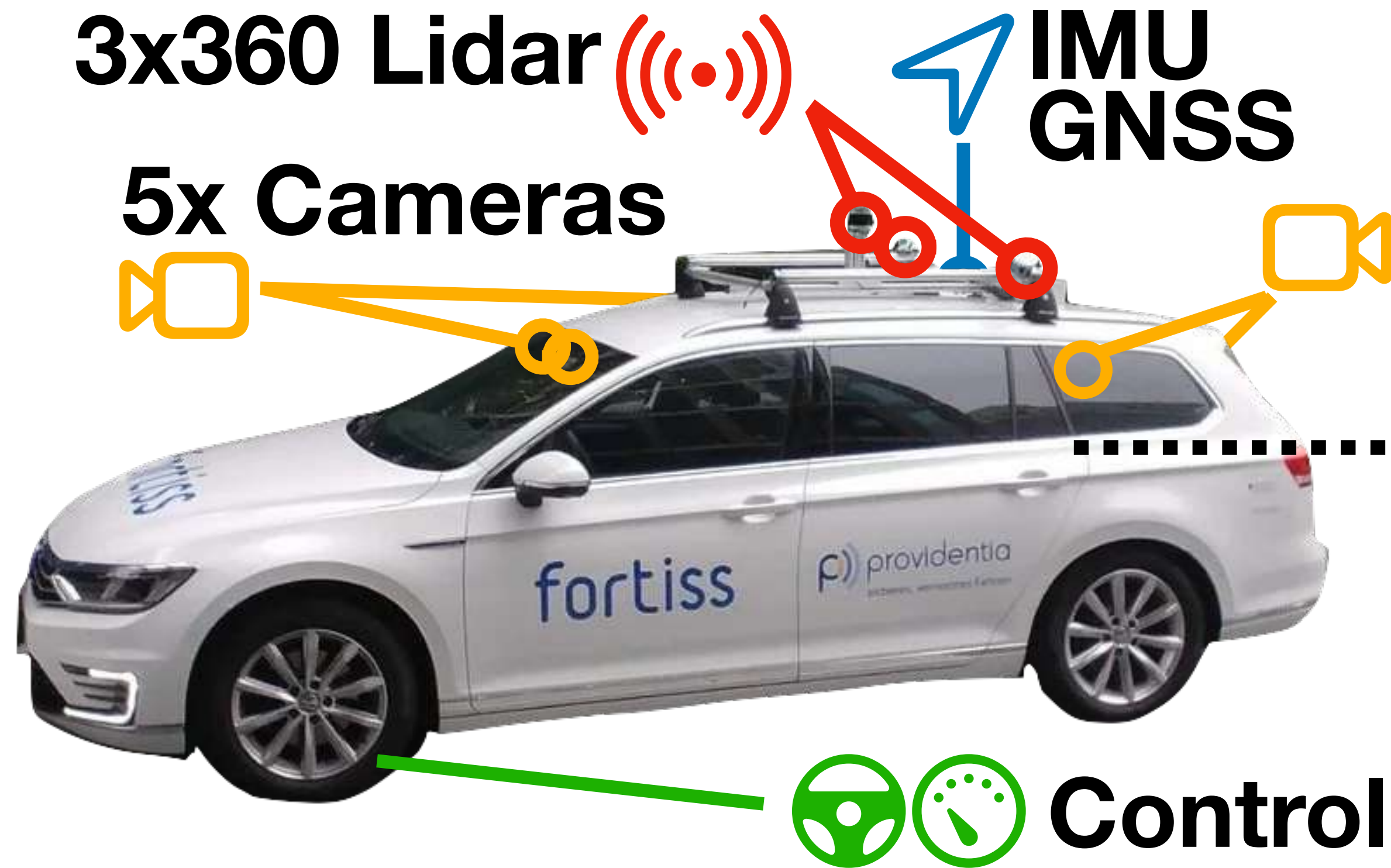
Full-scale Real-world testing

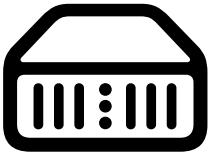
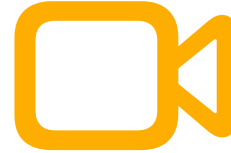
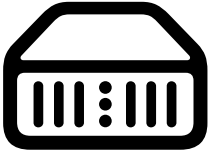
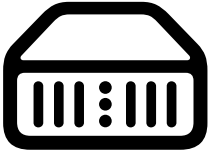
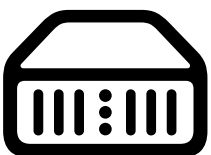


Real

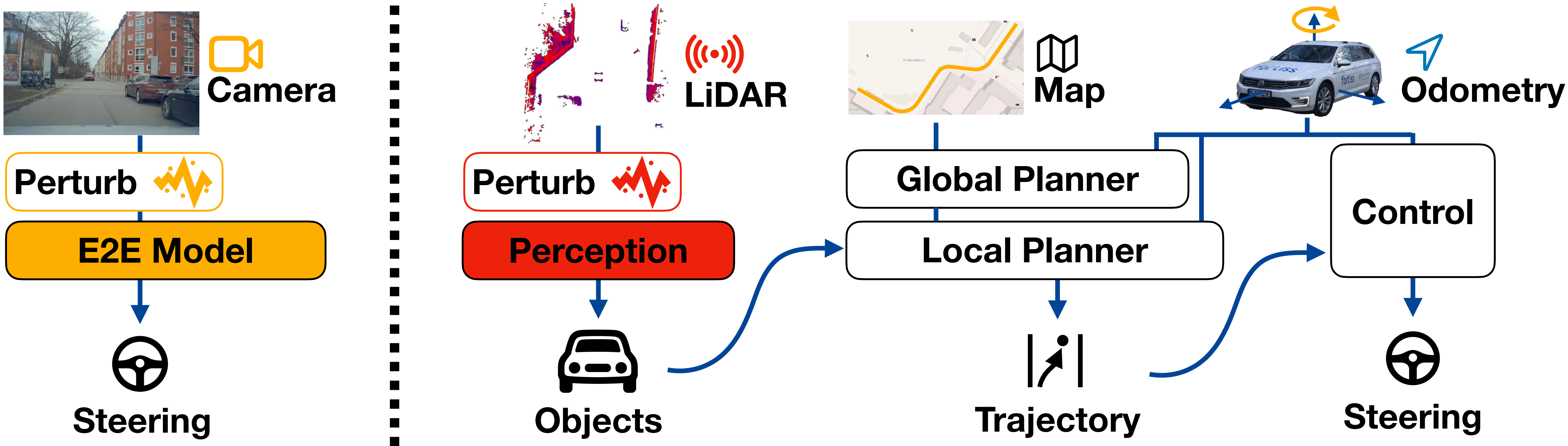
Behaviour
Scenarios
Sensors

Real World Perturbation Testing



	Volkswagen Passat GTE
	 Ultra7 - 64 GB - RTX5080
	 Nvidia Drive PX 2
	 i7 - 32 GB - GTX1050 Ti
ROS	 i7 - 32 GB - GTX1050 Ti
 	 dSpace Micro Autobox II

Real World Perturbation Testing





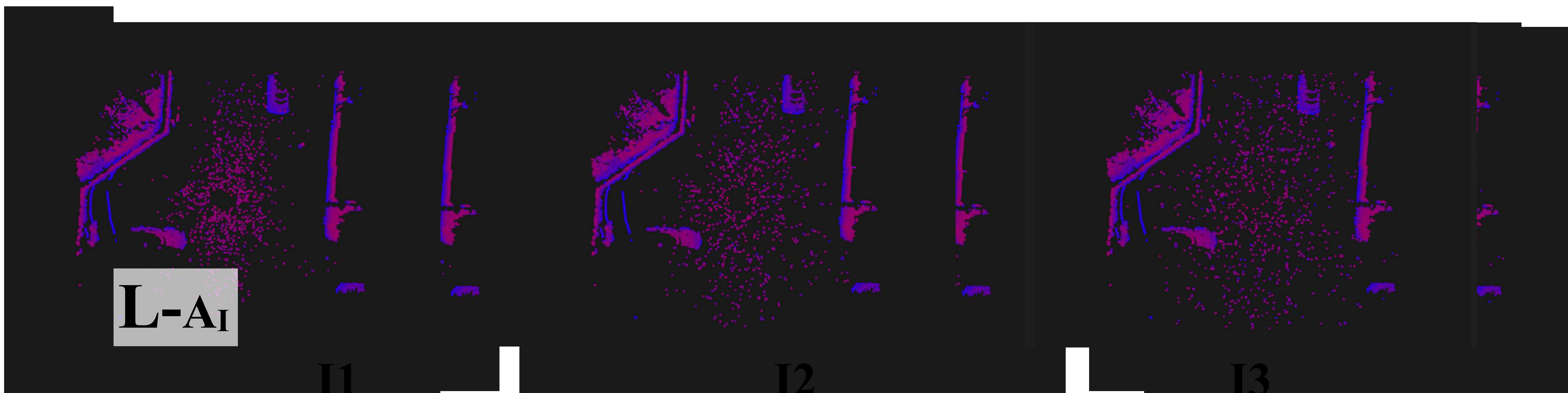
I1

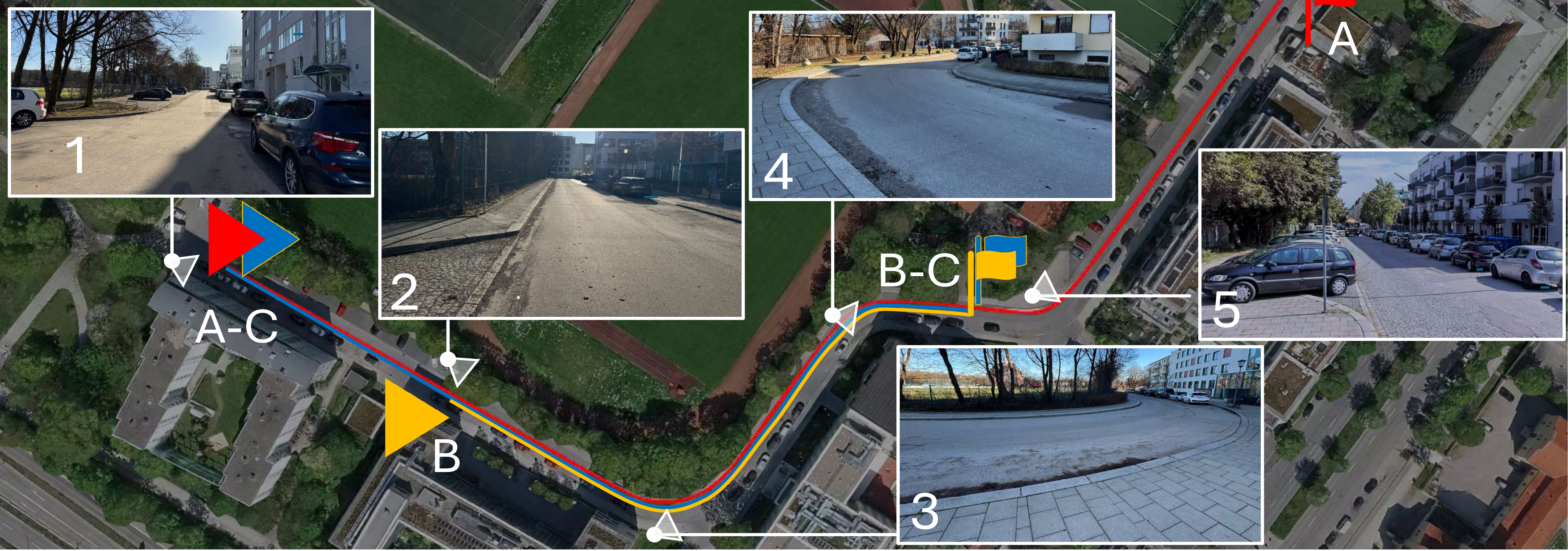


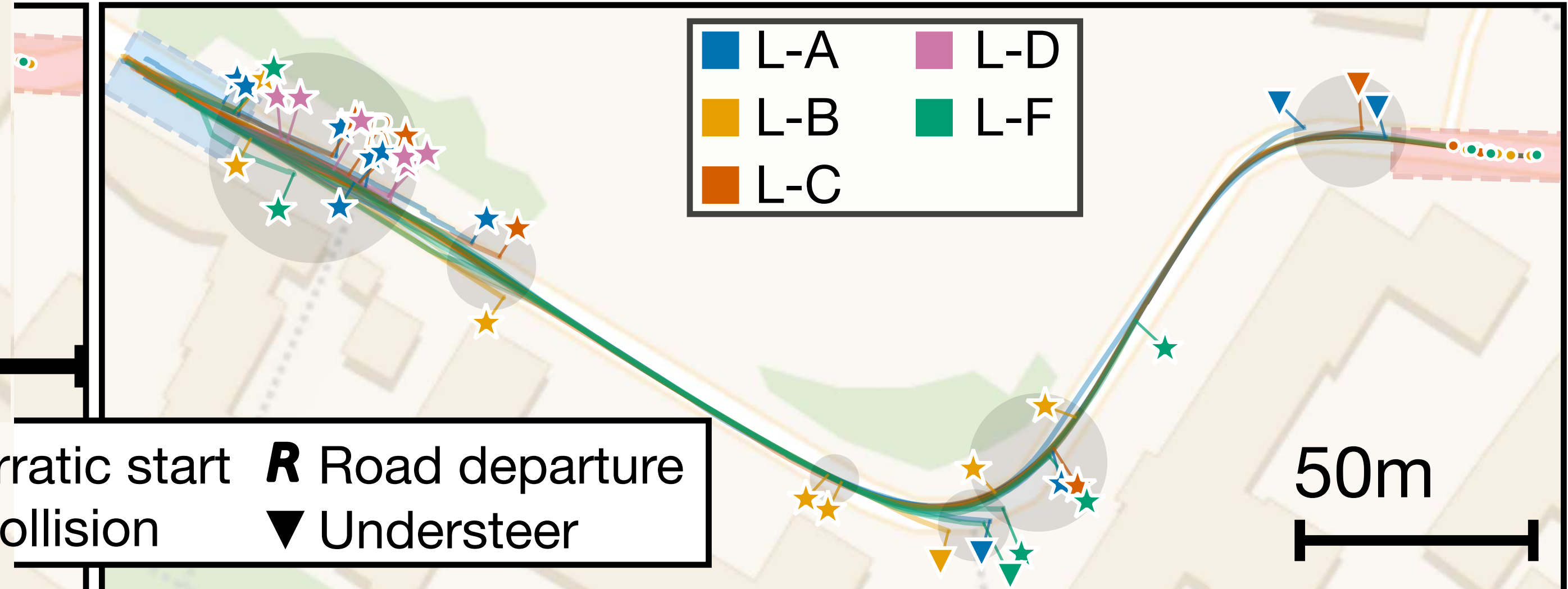
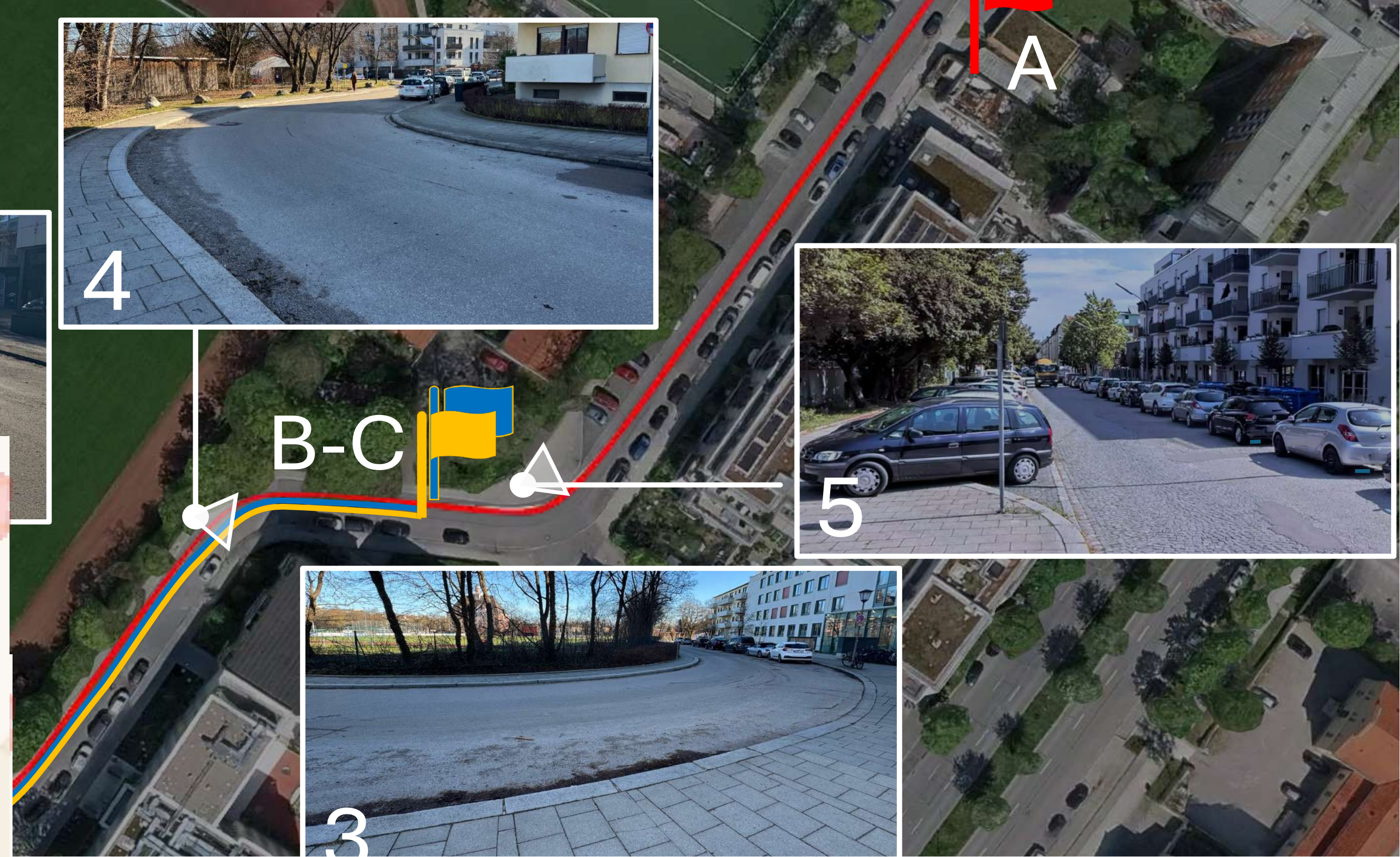
I3



I5







Failures

Oversteer



Erratic start

Collision



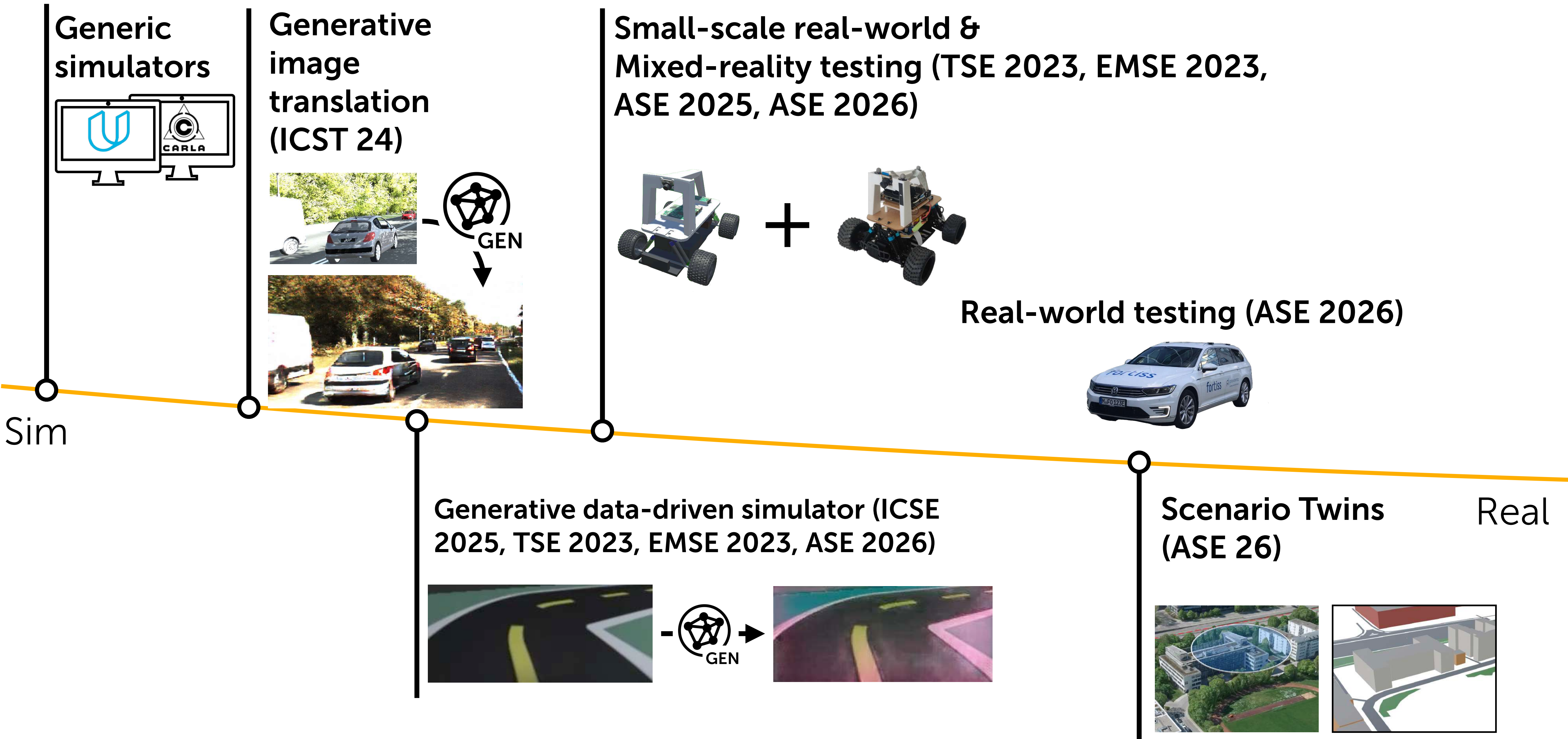
Road departure



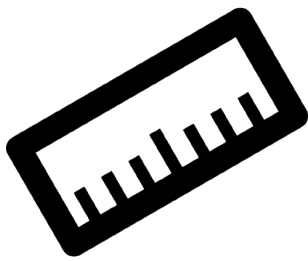
Understeer

50m

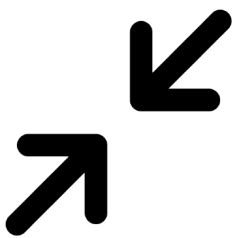
Reality Gap Mitigation



Our techniques

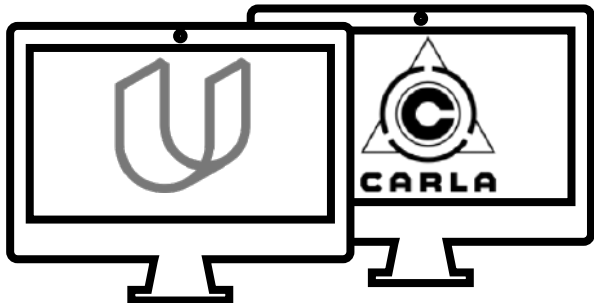


Measure

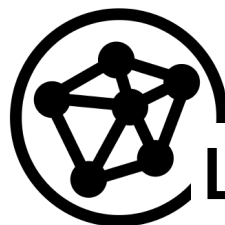


Mitigate

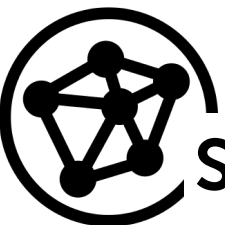
Generic
simulators



Sensors
Physics
Scenarios



LLaVA



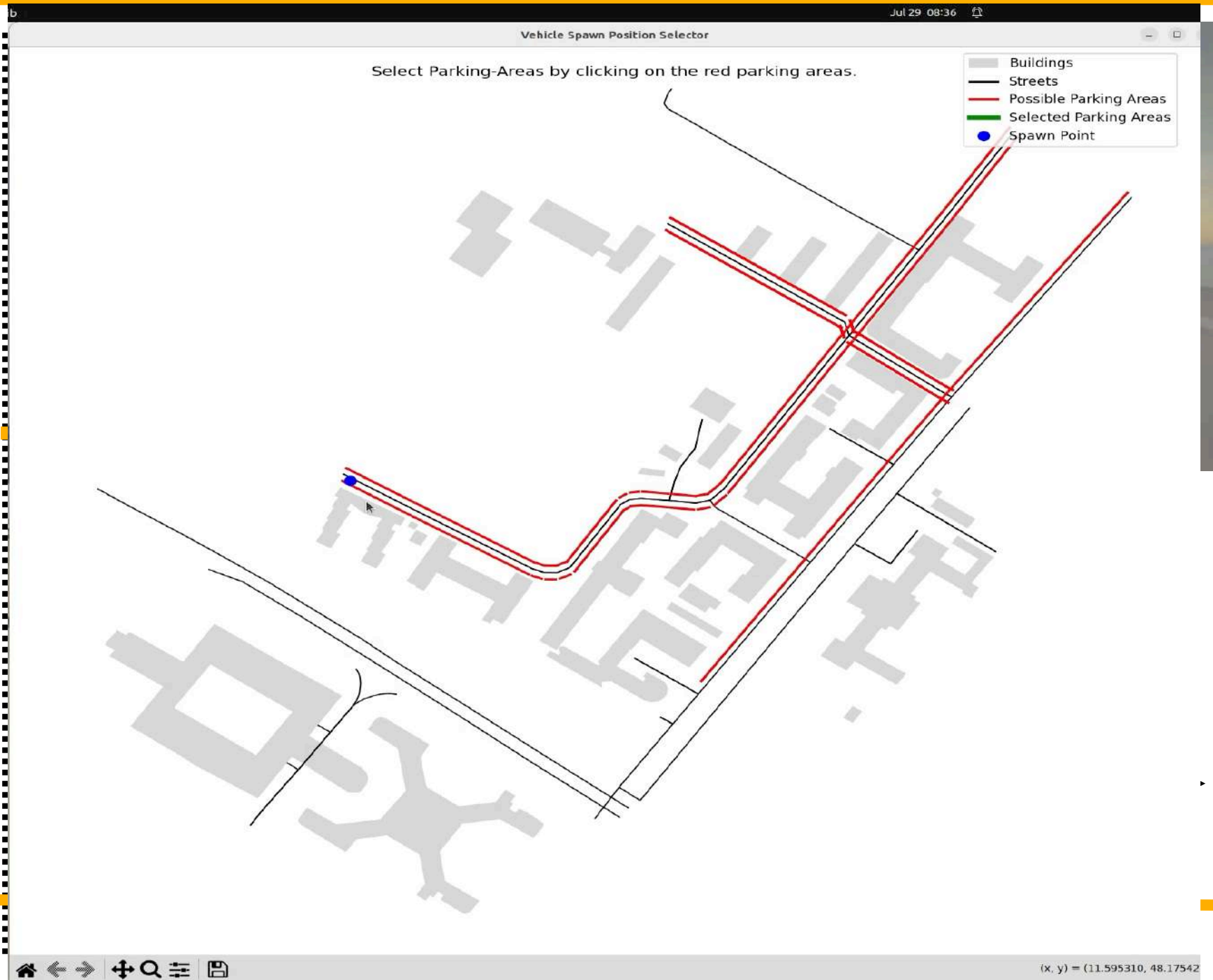
SegFormer

Dataset



Open
Street
Map

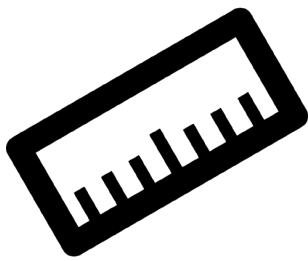
.osm



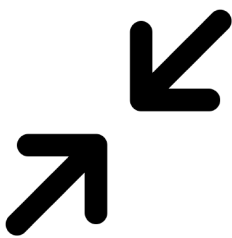
Sim

Real

Our techniques

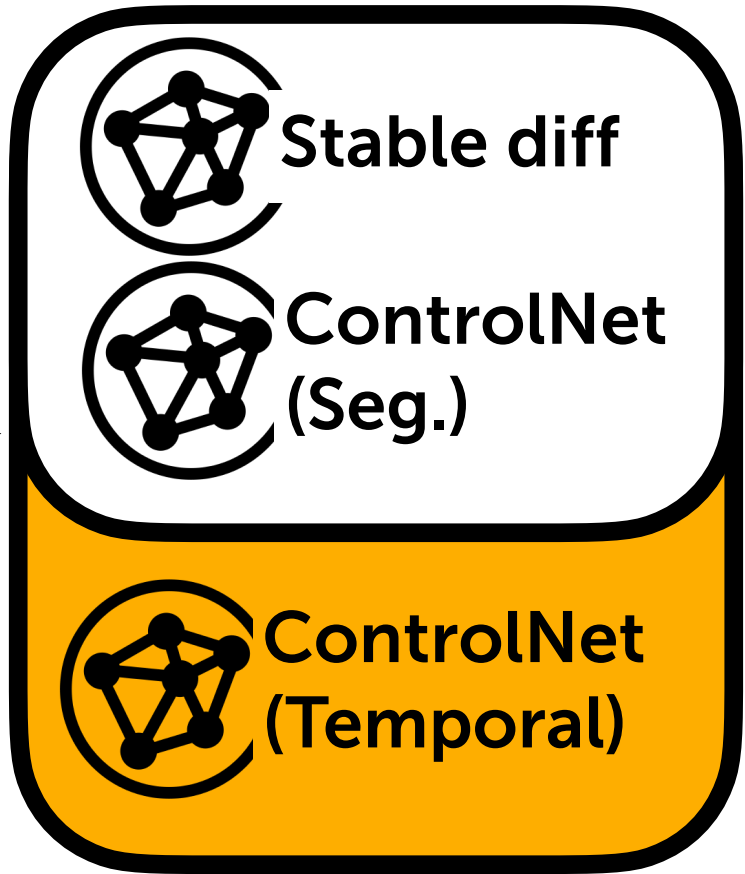
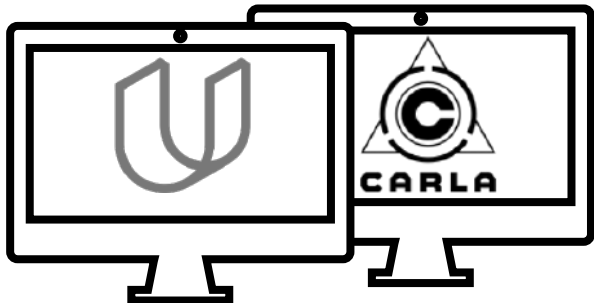


Measure



Mitigate

Generic
simulators



Sim

Real

Neural Reconstruction

From 2D datasets to 3D simulation



Neural Reconstruction

From 2D datasets to 3D simulation



Real-world Testbeds

From prototype to the real world



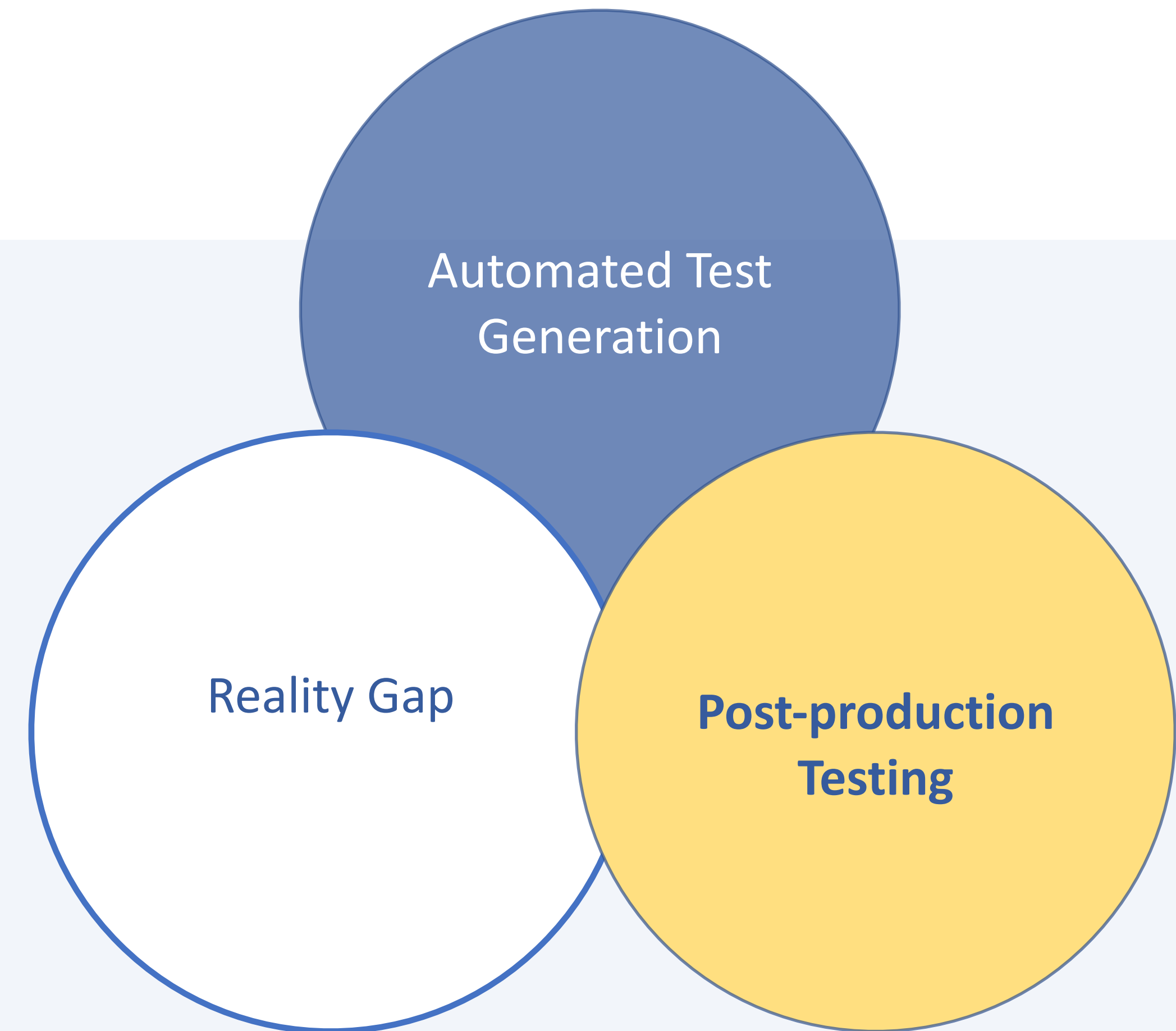
Takeaways

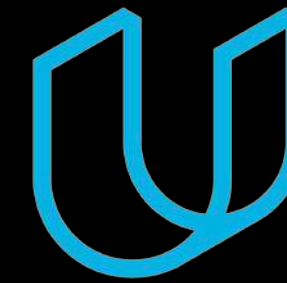
- ① **Reality gap affects scenarios, sensors, and physics**
- ② **Tackle one at a time (e.g., genAI for sensors). Metrics!**
- ③ **World Models, Physical AI?**

Automated Software Testing

Main research topics

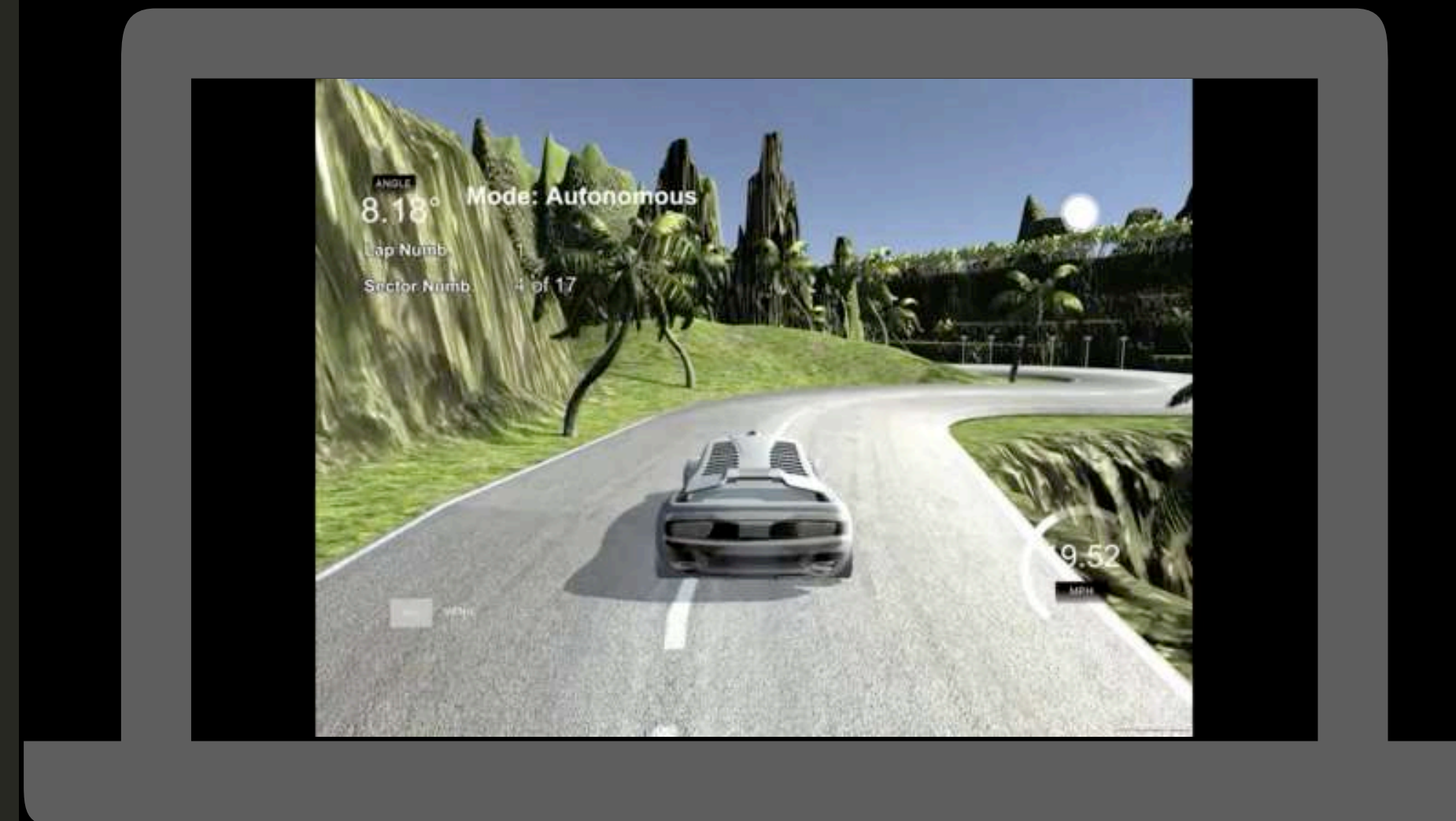
- **Automated Test Generation**
How can we automatically generate complex scenario-based tests efficiently and effectively?
- **Reality Gap**
How can we measure and mitigate the reality gap between simulation and in-field testing?
- **Post-production Testing**
How to ensure a high dependability of deep neural network driven-cyber-physical systems (CPS) in production?





Steering Angle predicted by the DNN

```
-0.16933852434158325  
-0.2953818142414093  
-0.2953818142414093  
-0.2953818142414093  
-0.24482464790344238  
-0.24482464790344238  
-0.24482464790344238  
-0.2340604066848755  
-0.2340604066848755  
-0.2340604066848755  
-0.2876757085323334  
-0.2876757085323334  
-0.2876757085323334  
-0.28597092628479004  
-0.28597092628479004  
-0.28597092628479004  
-0.280177503824234  
-0.280177503824234  
-0.280177503824234  
-0.1850987821817398  
-0.1850987821817398  
-0.1850987821817398  
-0.2626234292984009  
-0.2626234292984009  
-0.2626234292984009  
-0.20239685475826263
```



* vehicle learnt how to stay in lane

**Training set cannot contain
ALL possible driving
conditions!**

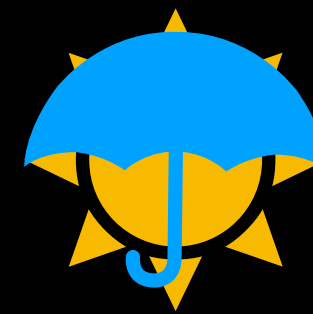
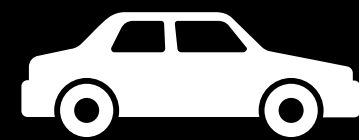
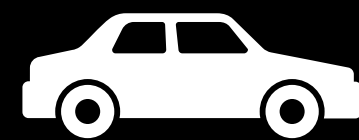
A hand is holding a clear crystal ball. Inside the crystal ball, the text "Can we predict unexpected conditions" is written in a black, italicized serif font. A small red asterisk is placed after the word "conditions". To the right of the crystal ball, a large, semi-transparent grey question mark is visible.

*Can we **predict**
unexpected conditions**

** Unexpected condition:*

*Unseen, potentially hazardous and
misbehaviour-inducing driving
condition*

Unexpectedness
metric



unexpectedness

prediction

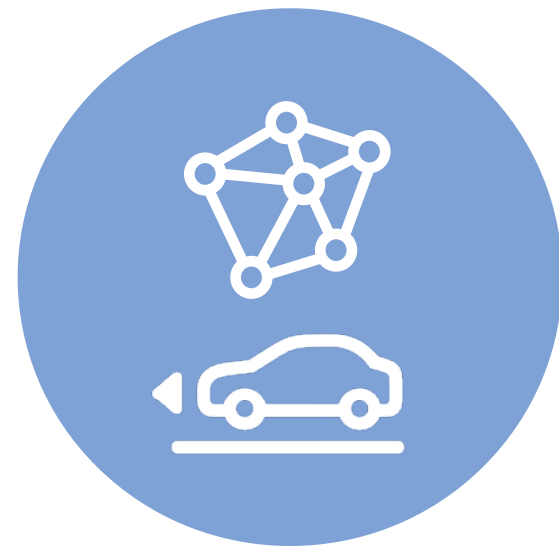


threshold

time

Post-production Testing

Real-time monitoring of deployed ADS



ADS & ADAS

- End-to-end DNNs (level 2)
- Full AD stacks



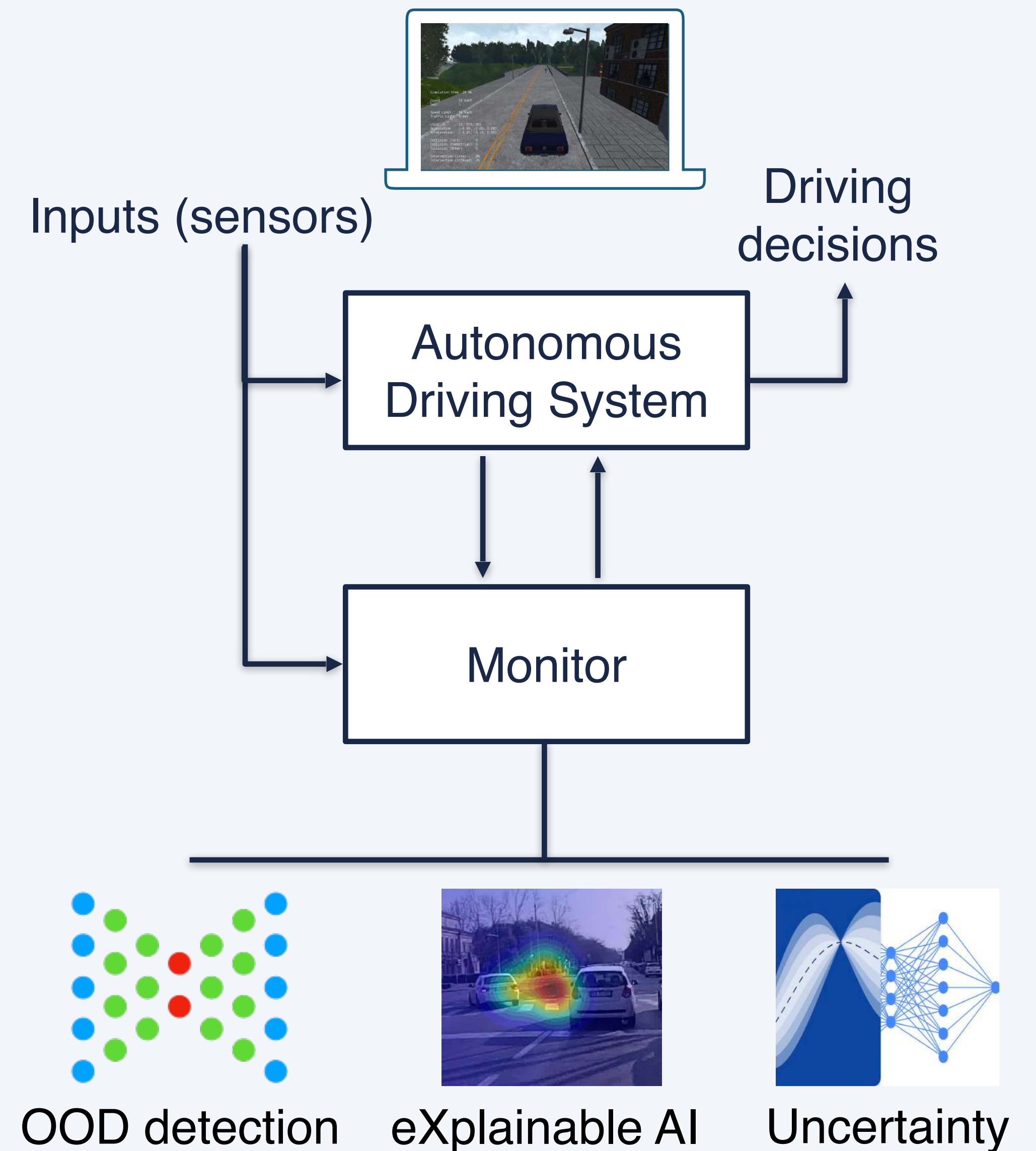
Real-time Monitoring

Observe System Under Test at runtime

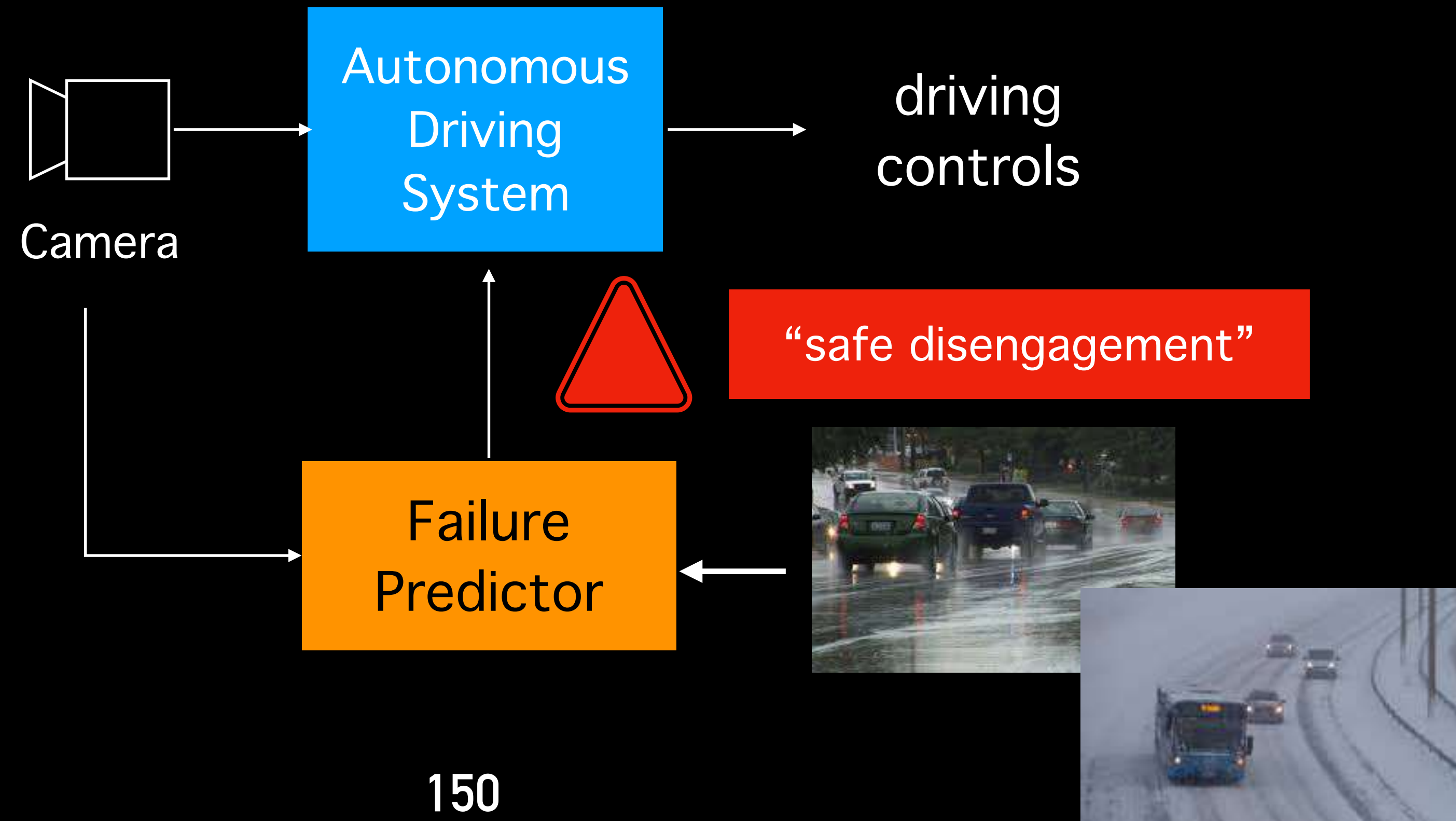


Trustworthy ADS

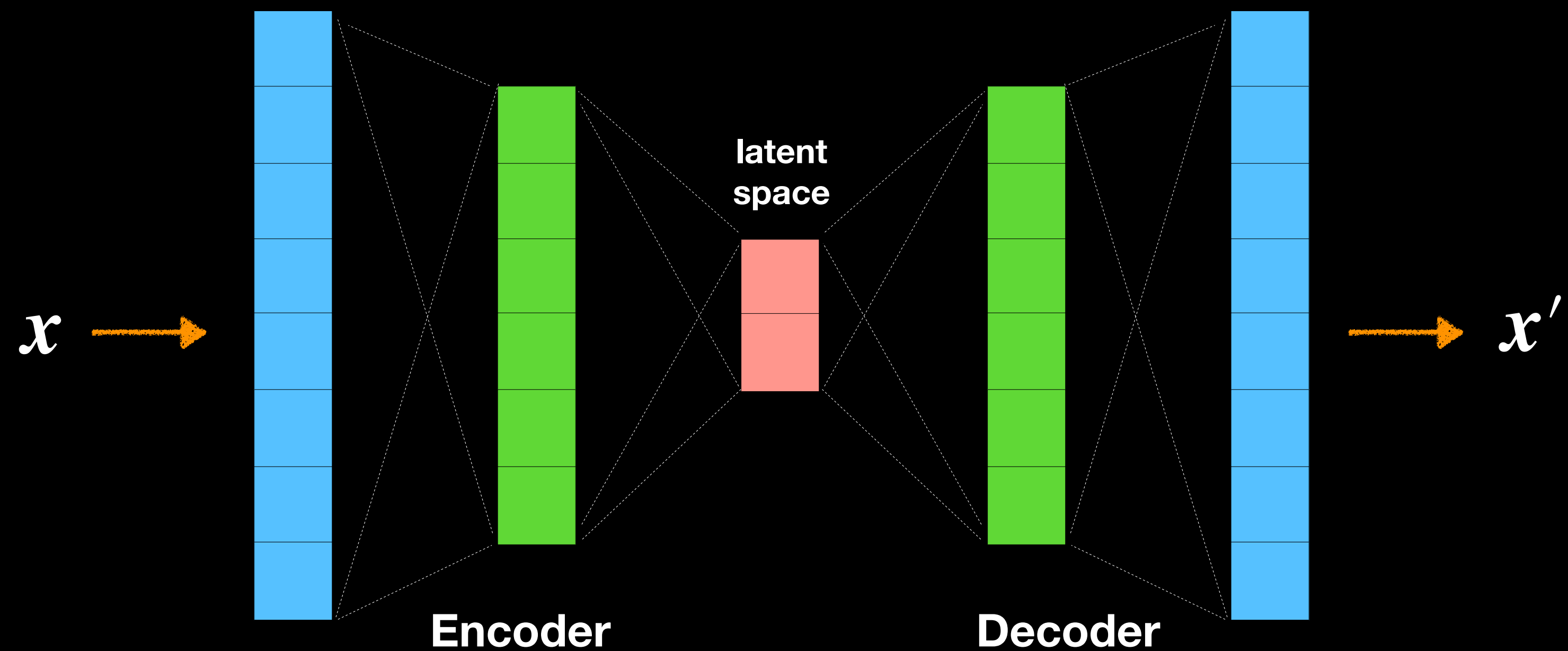
Alert the driver if system is not to be trusted or activate fallback procedures



Black Box Unexpectedness Score



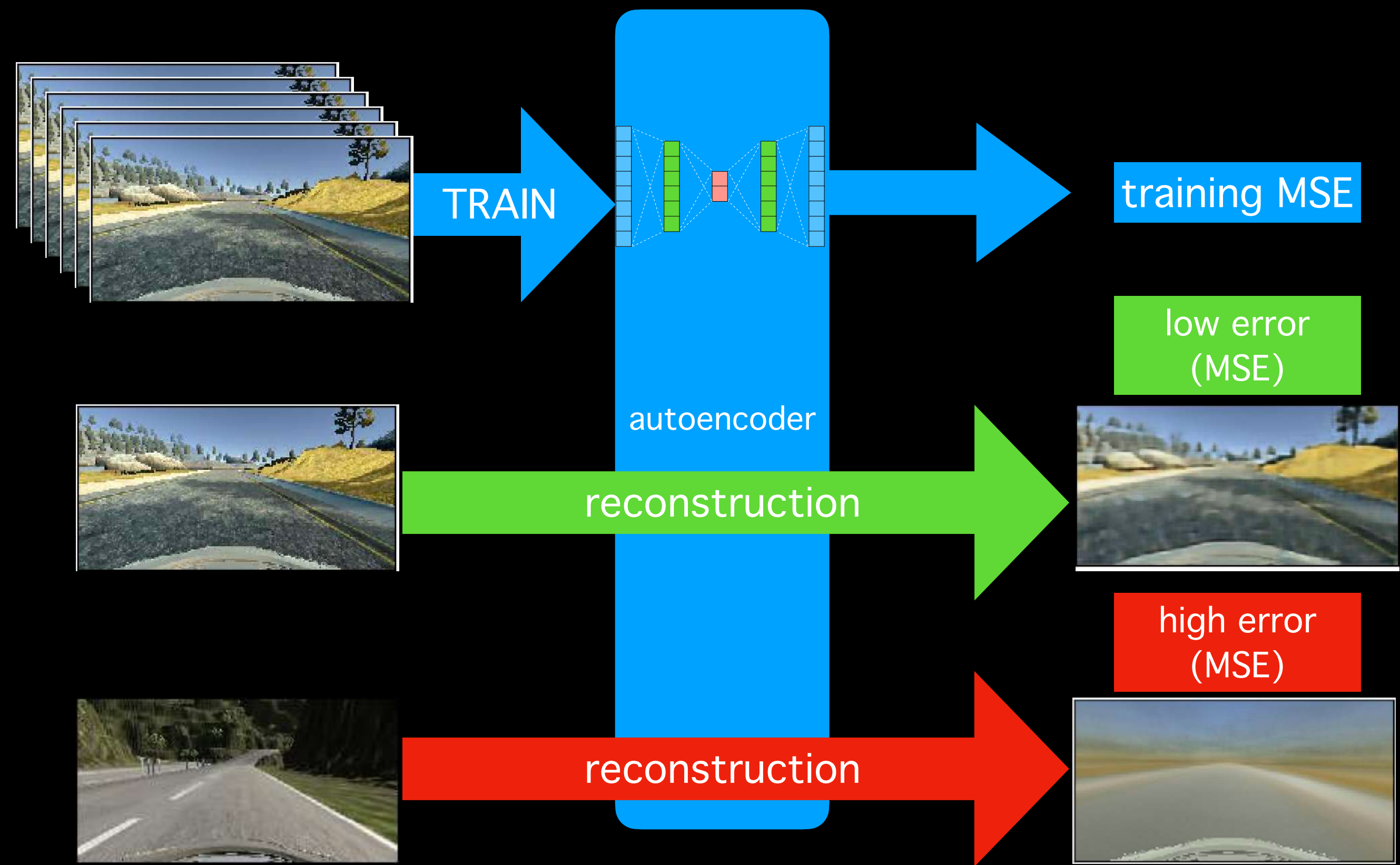
Autoencoders



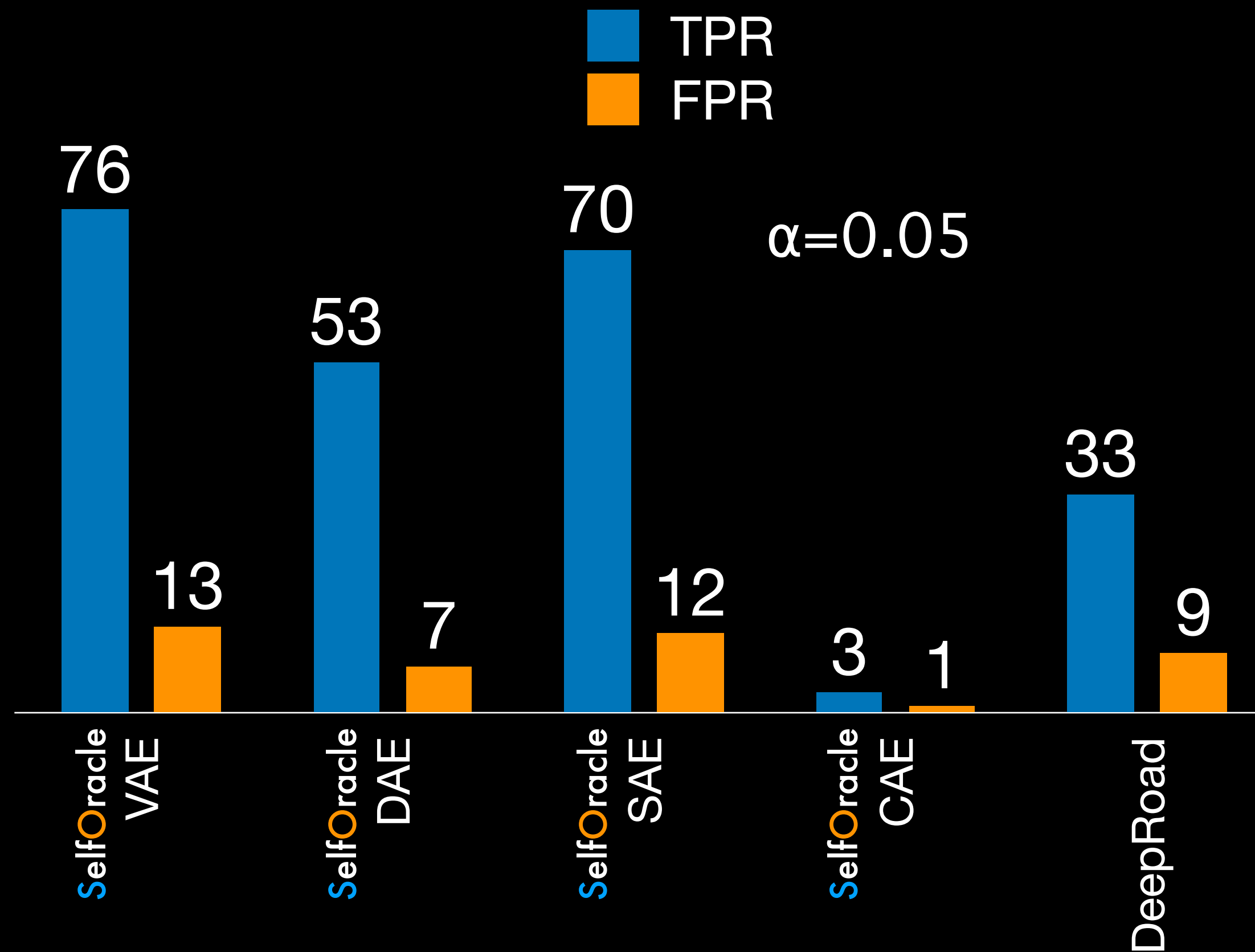
$$MSE(x, x') = \frac{1}{n} \sum_{t=1}^n \|x - x'\|^2$$

indicator of the difference
between the original data and
its reconstruction

Autoencoder-based Anomaly Detector

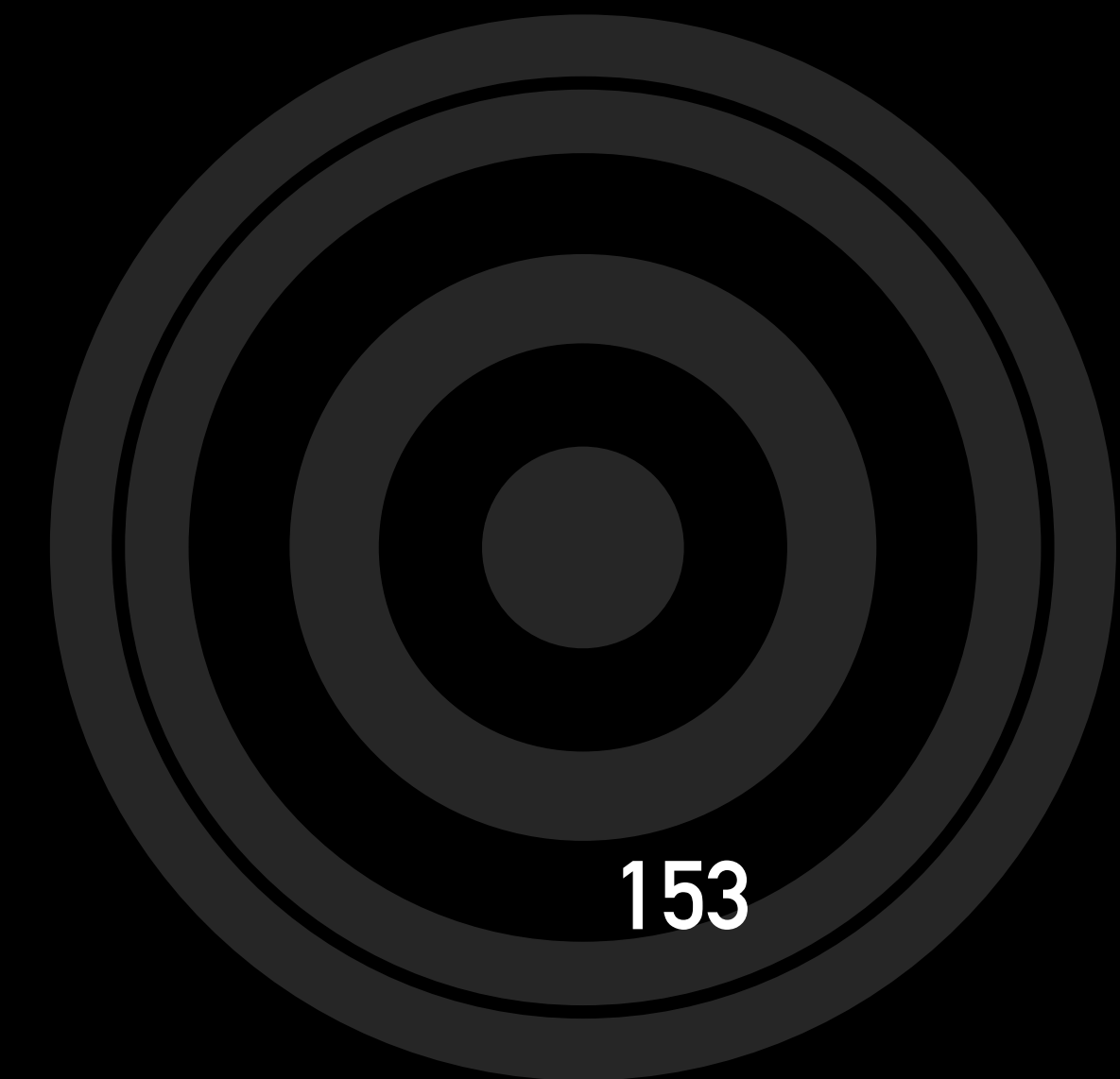


Is our approach **effective**
in **predicting** misbehaviours ?

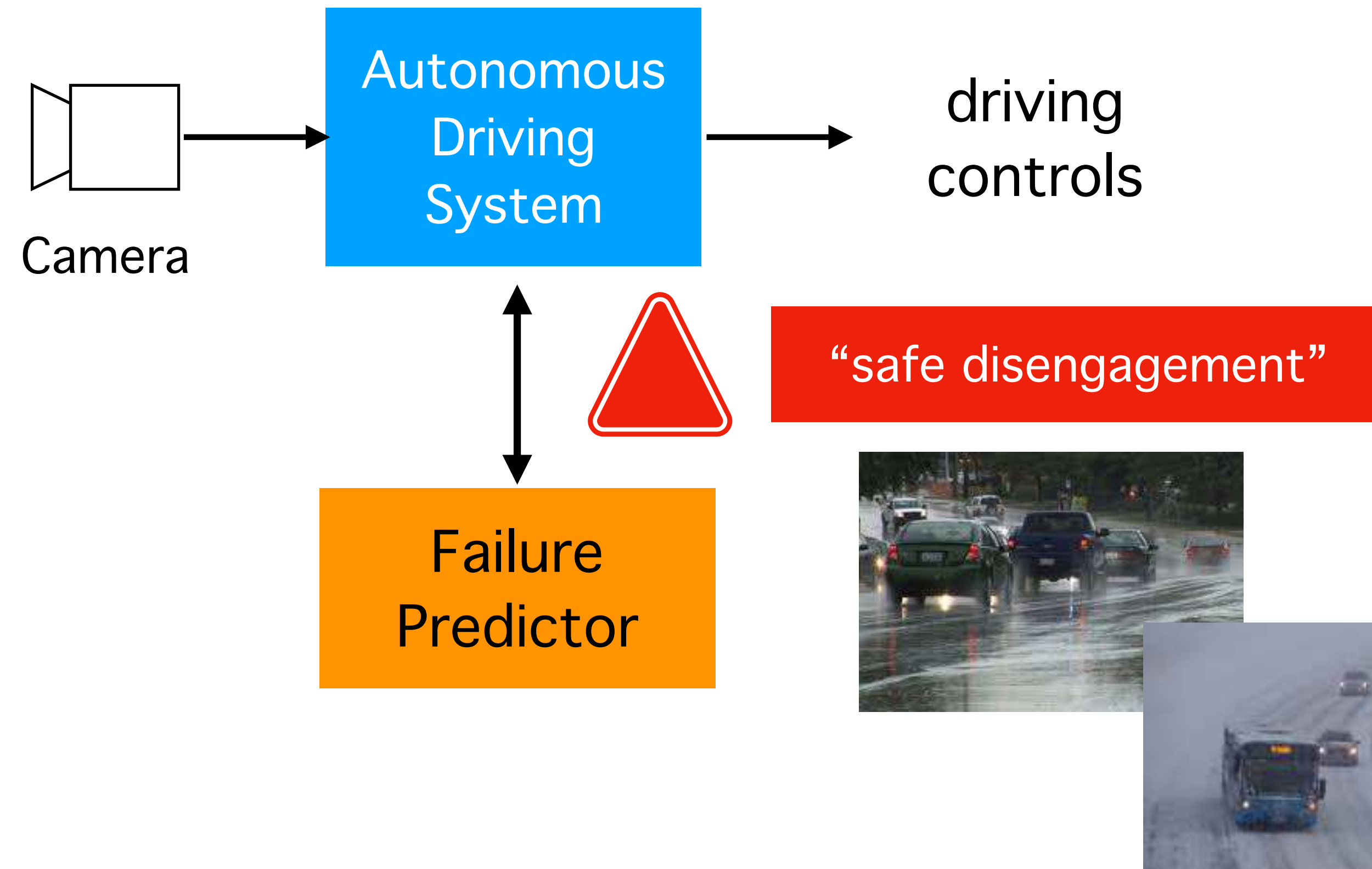


76%

correct misbehaviour
predictions **6 to 9**
seconds before



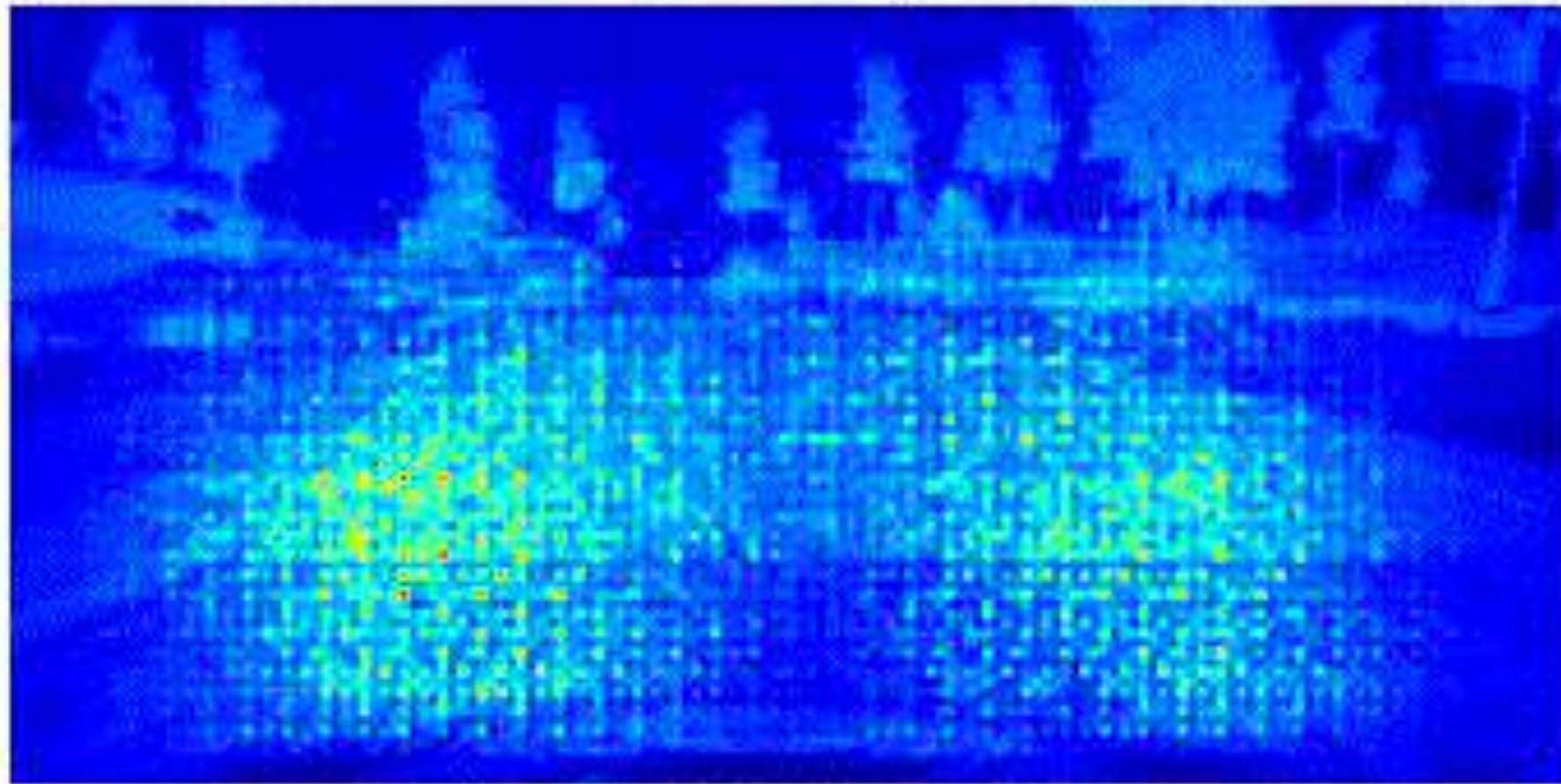
White Box Unexpectedness Score



XAI for failure prediction

Nominal

XAI Image



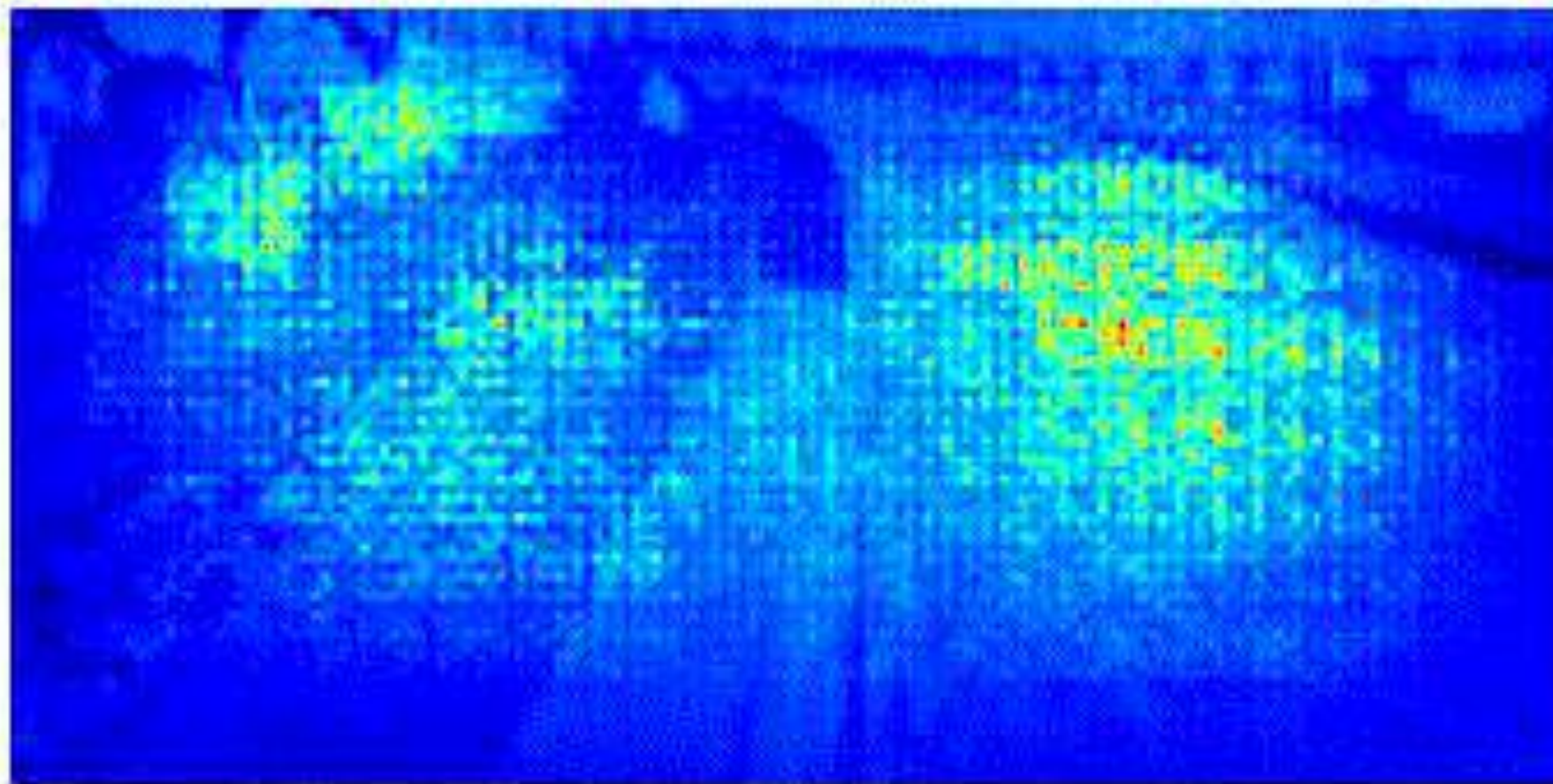
Original Image



XAI for failure prediction

Uncertain/Unexpected

XAI Image



Original Image



Confidence Score Synthesis

$$\overline{h} = \frac{1}{WHC} \sum_{i=1, j=1, c=1}^{W, H, C} h_{[i][j][c]}$$

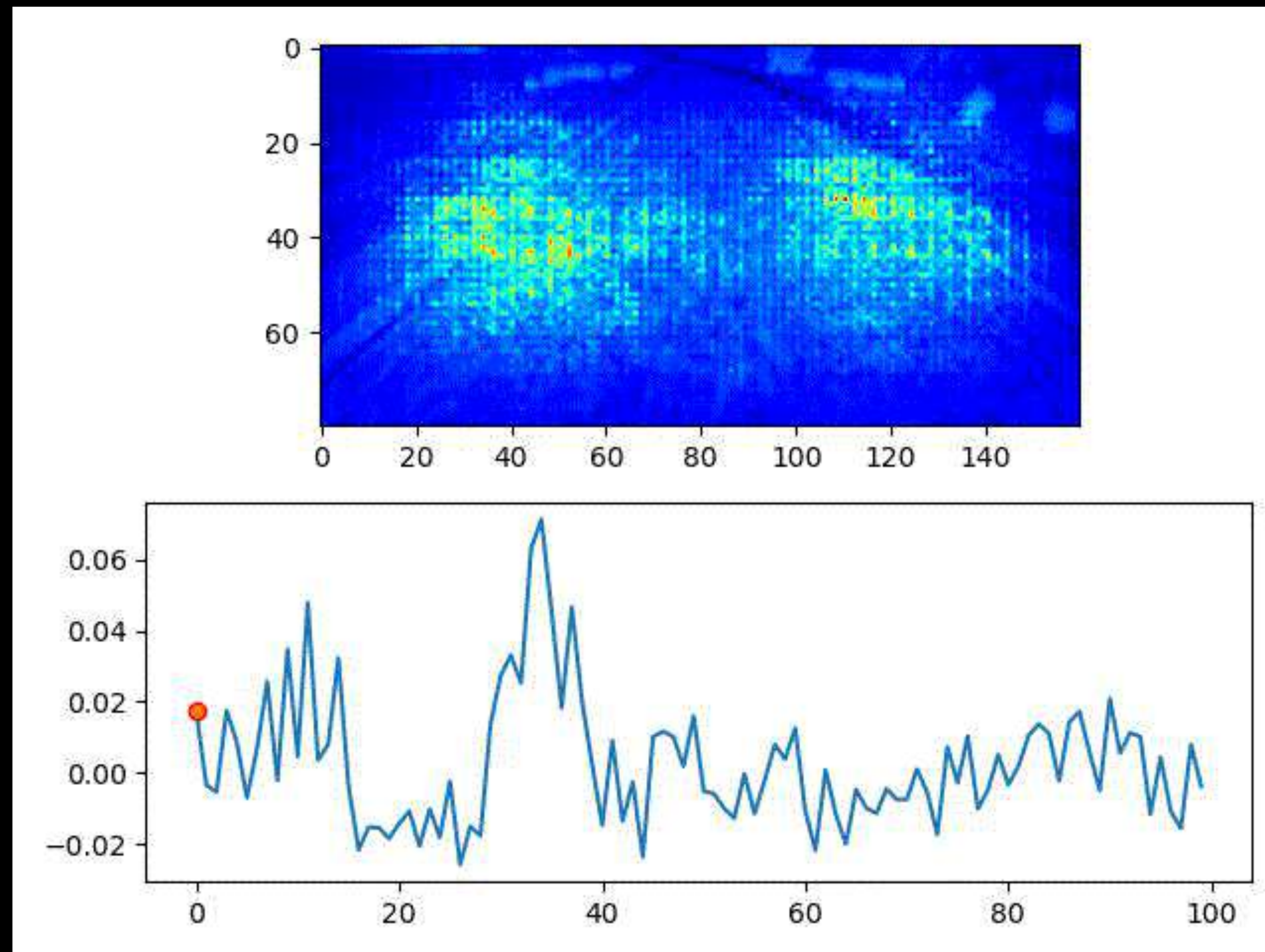
HA = Heatmap
Average Function

$$\overline{\nabla h_t} = \frac{1}{WHC} \sum_{i=1, j=1, c=1}^{W, H, C} h_{t[i][j][c]} - h_{t-1[i][j][c]}$$

HD = Heatmap
Derivative Function

$$h_e = \mathcal{L}(h, dec(enc(h)))$$

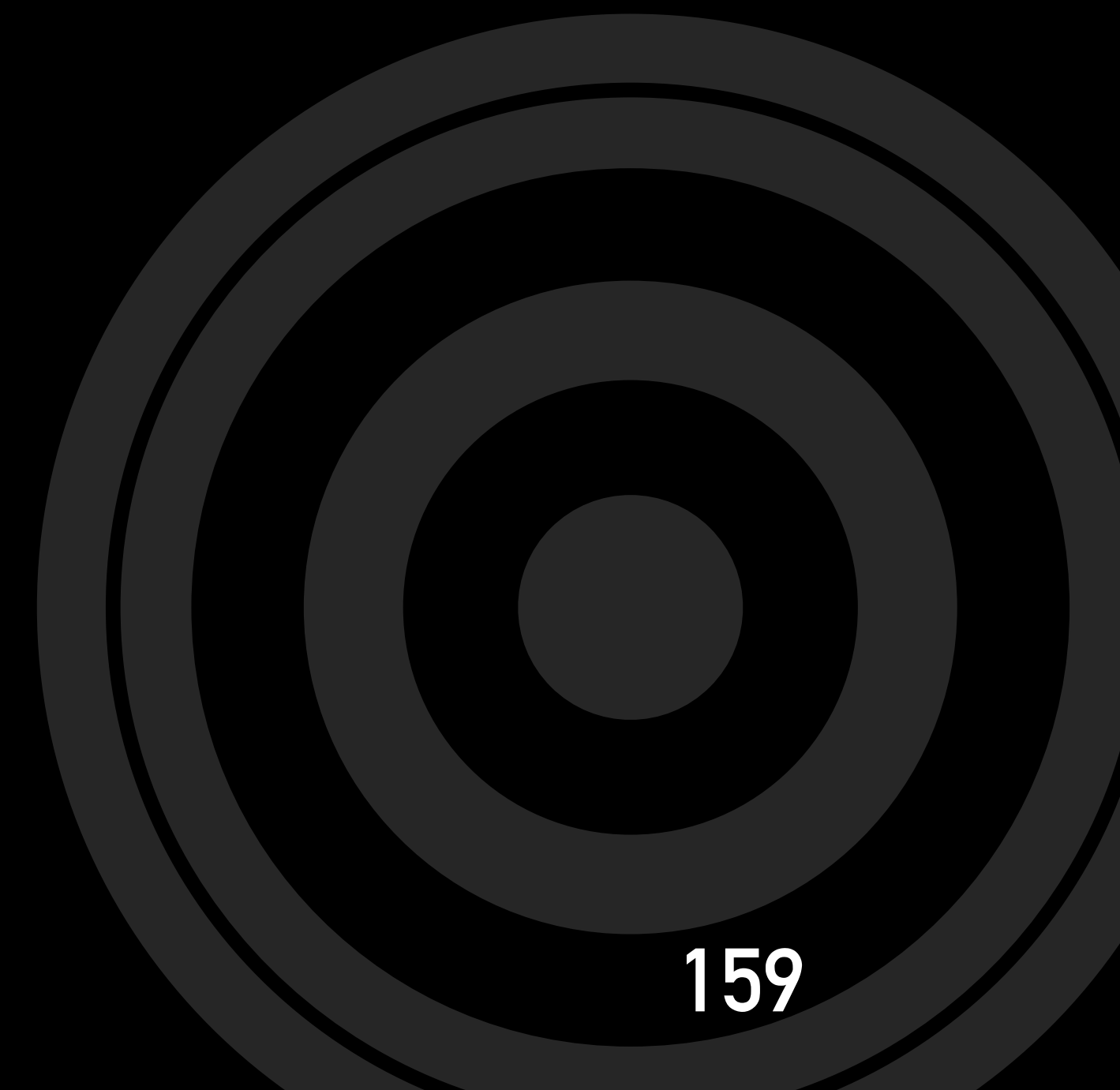
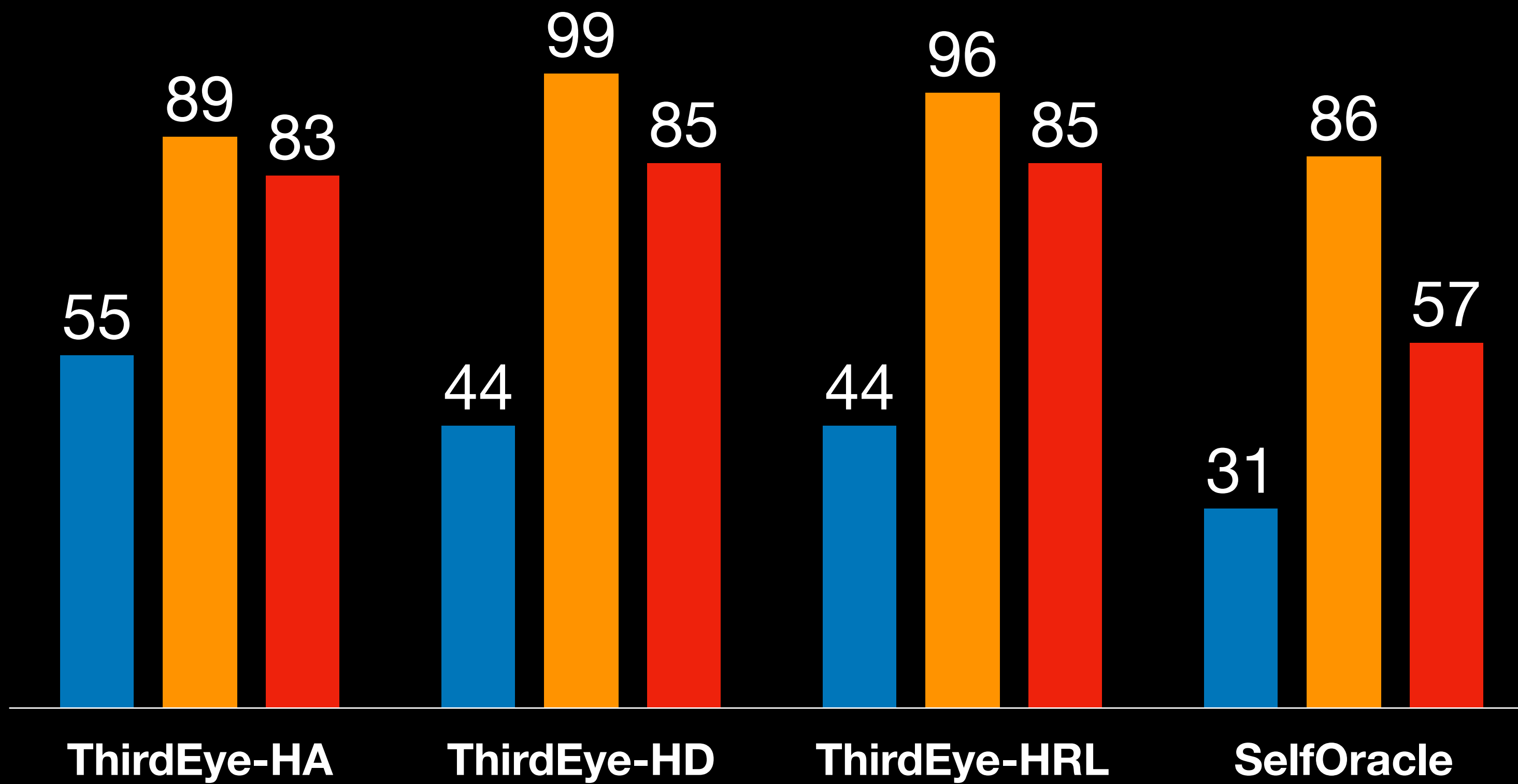
HRL = Heatmap
Reconstruction Loss



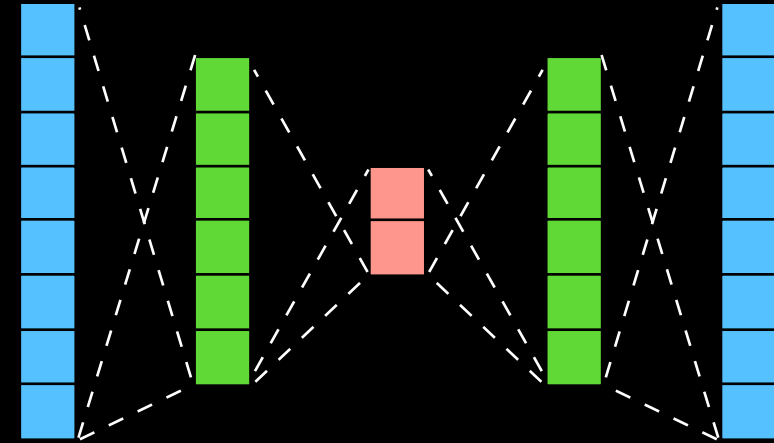
Is our approach **effective**
in **predicting** misbehaviours ?

Precision
Recall
F-3

$\alpha=0.05$



SelfOracle (ICSE 2020)

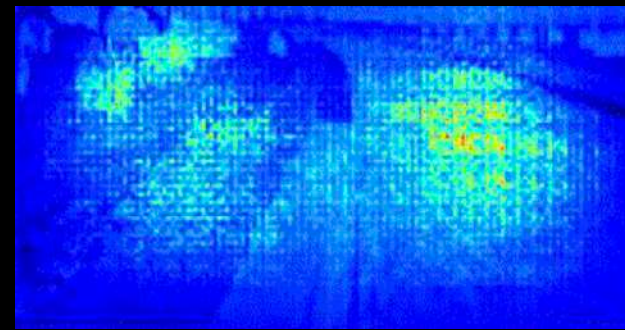
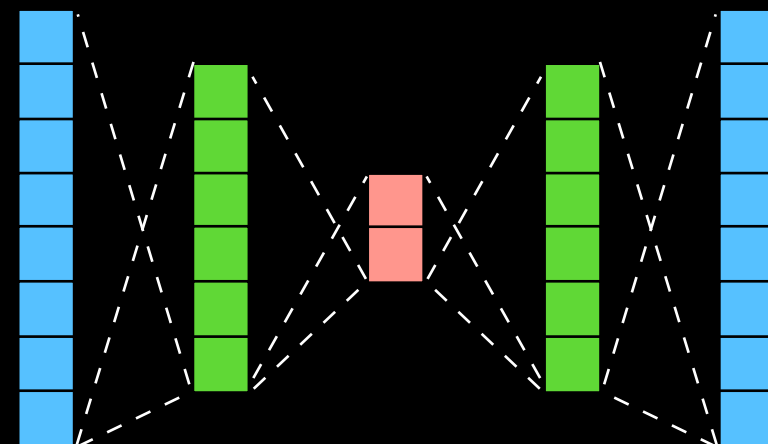
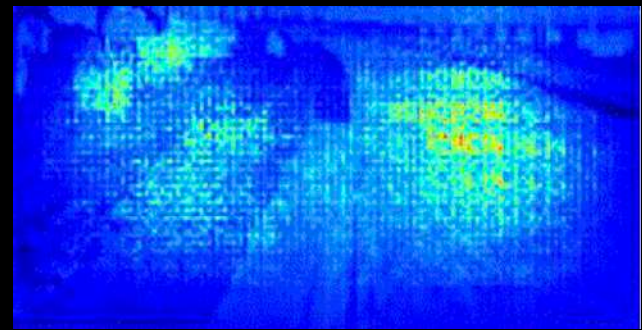


Autoencoder-based

Black-box

Computationally Fast

ThirdEye (ASE 2022)

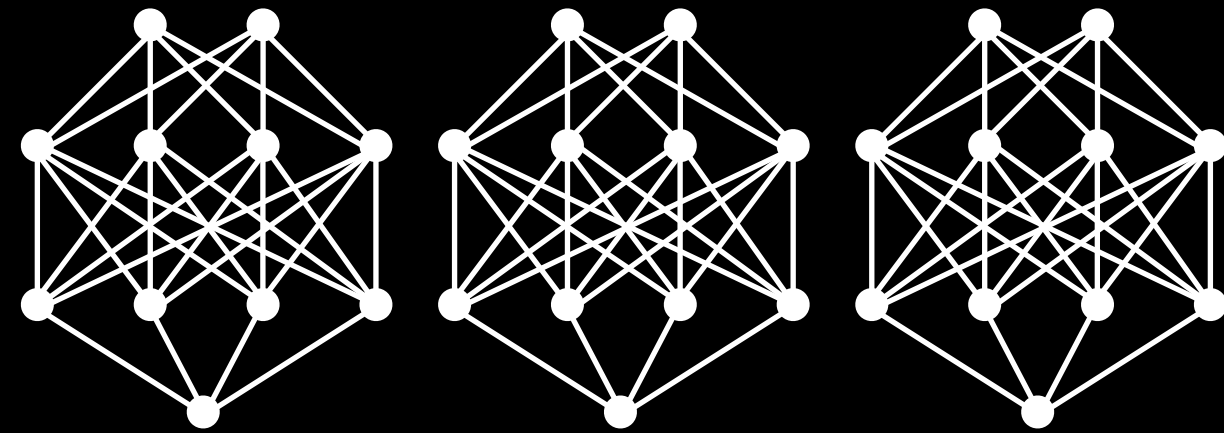


Heatmap-based

White-box

Computationally Expensive

Deep Ensembles



Uncertainty

score

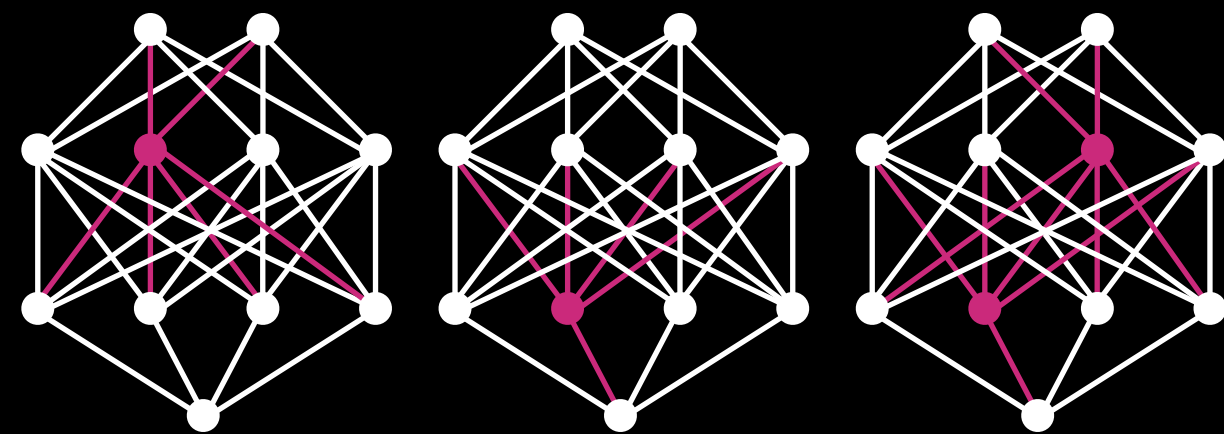


Theoretically grounded

Widely applicable

Computationally Expensive

Monte-Carlo Dropout



Uncertainty

score

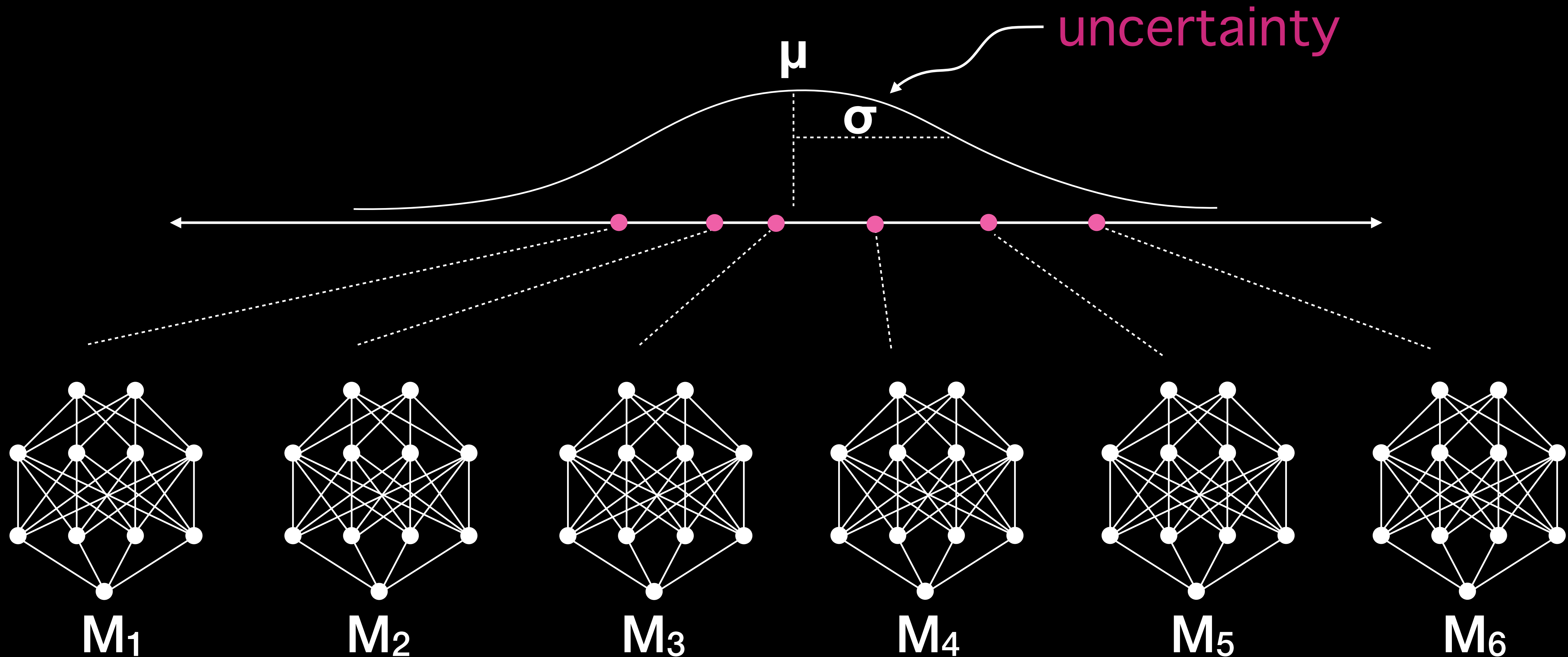


Approximation of BNNs

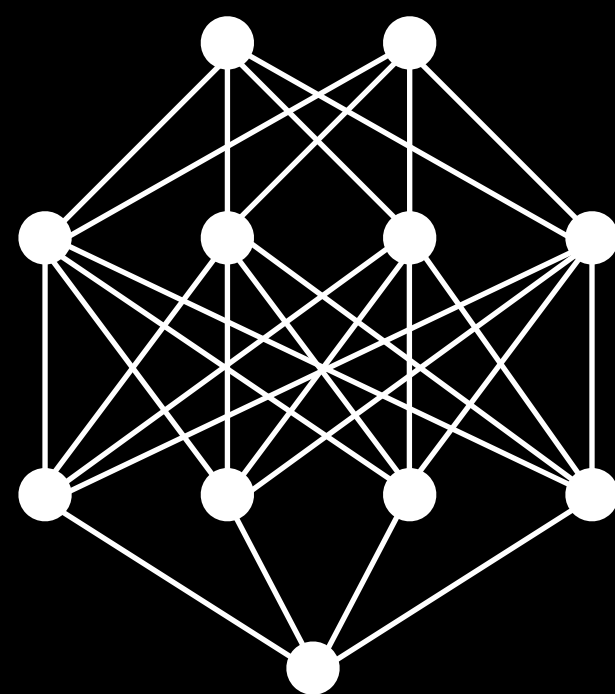
Intrusive Approach

Computationally Feasible

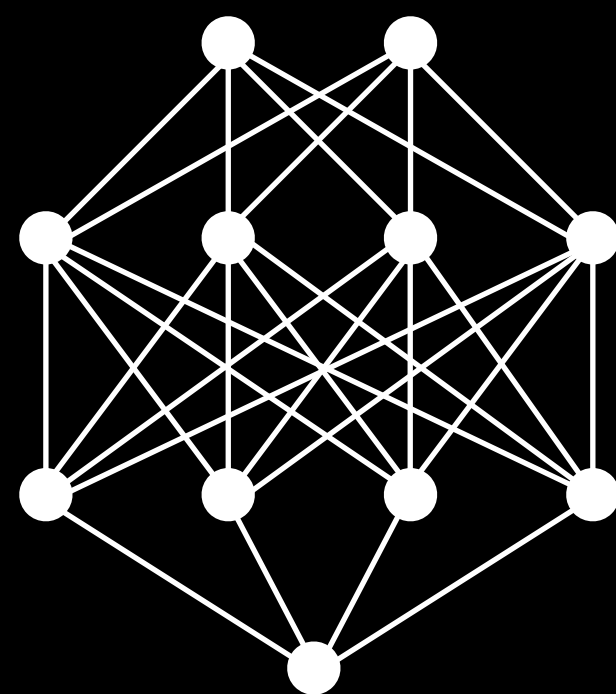
Deep Ensembles



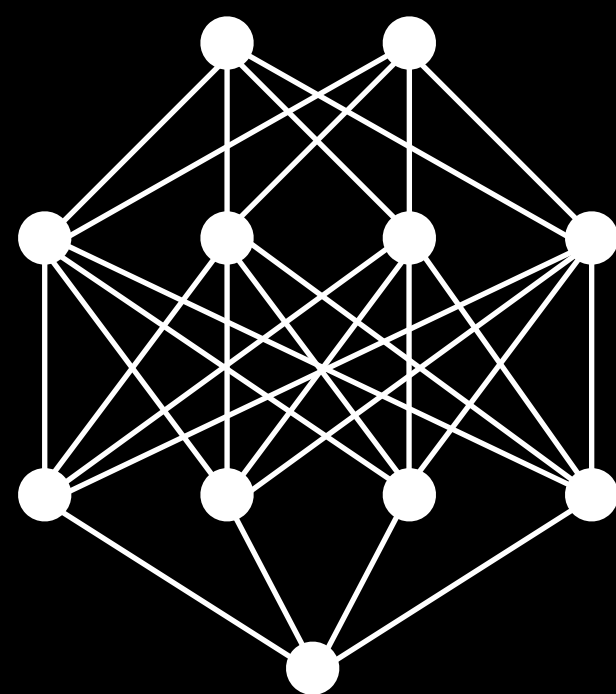
Monte-Carlo Dropout



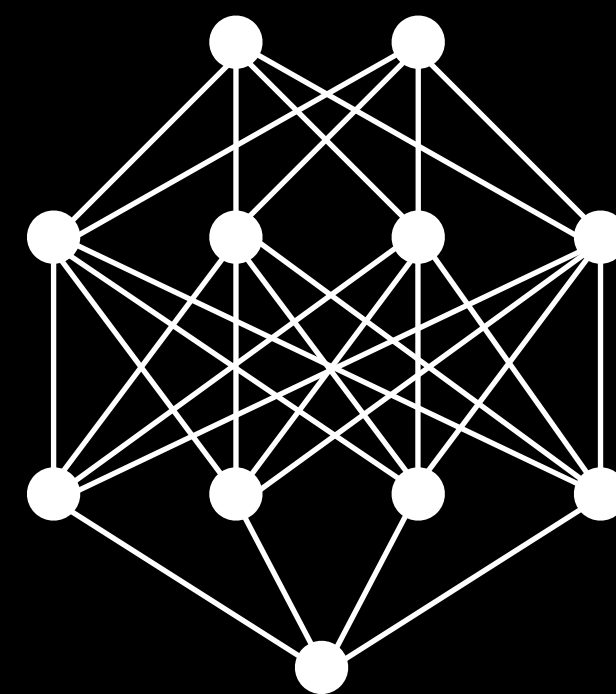
M_1



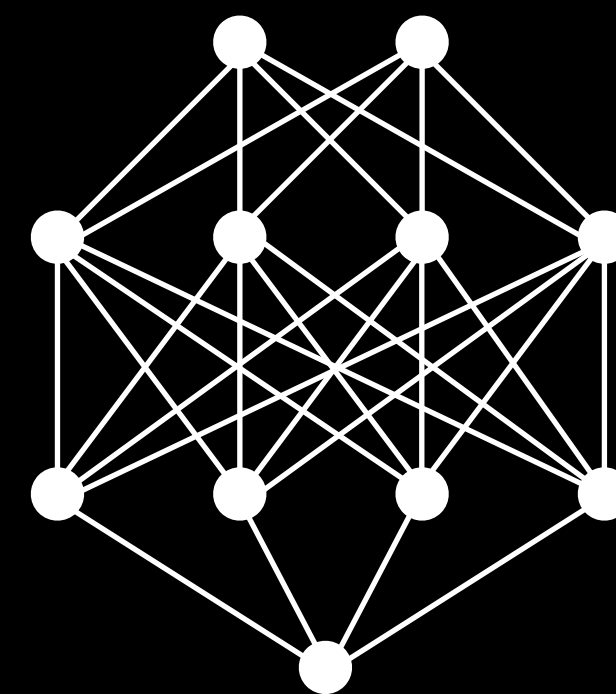
M_1



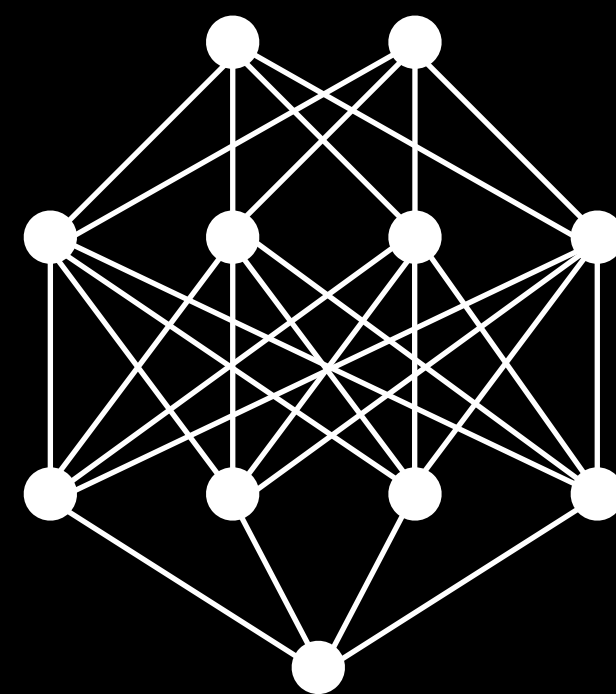
M_1



M_1

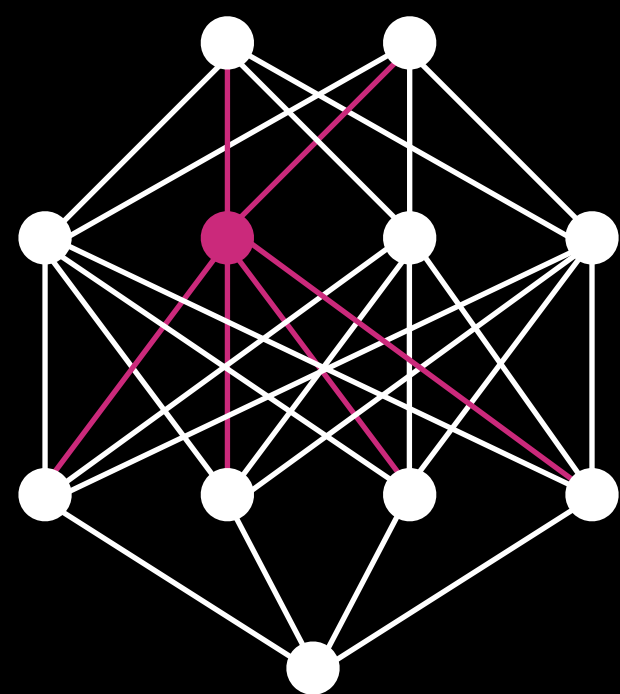


M_1

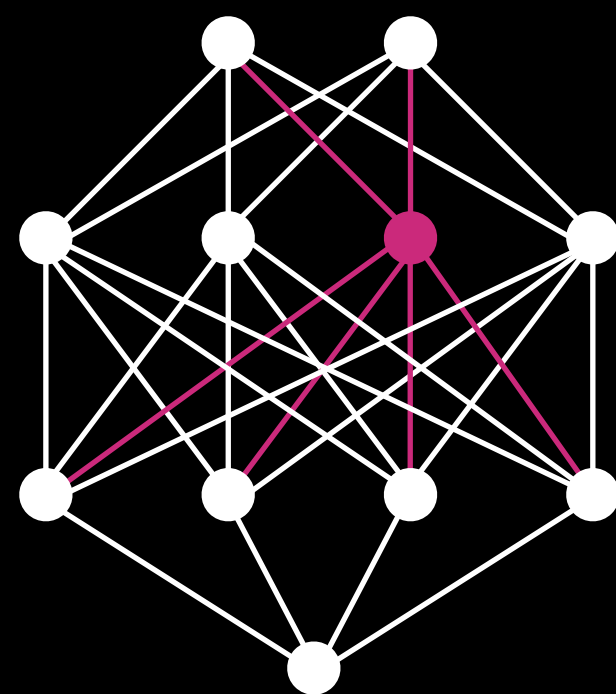


M_1

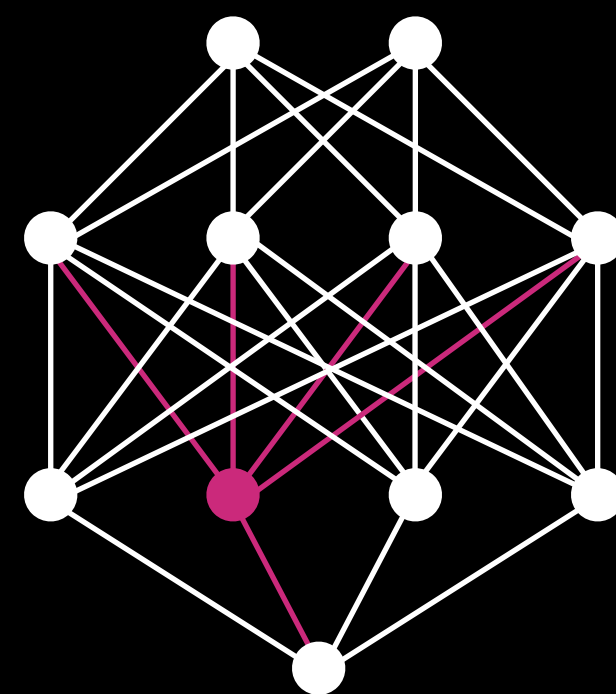
Monte-Carlo Dropout



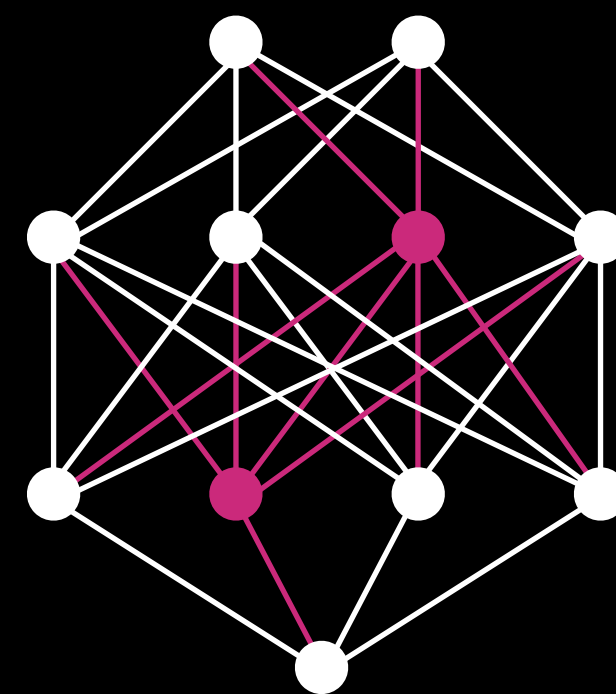
M_1



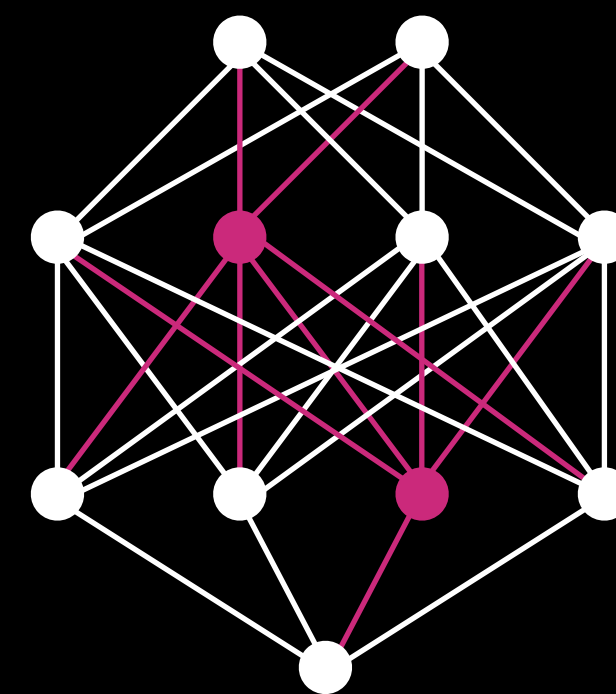
M_1



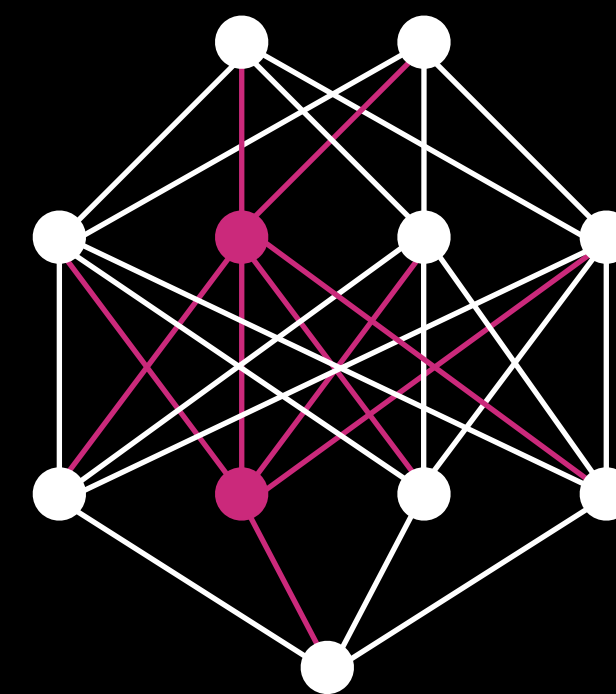
M_1



M_1

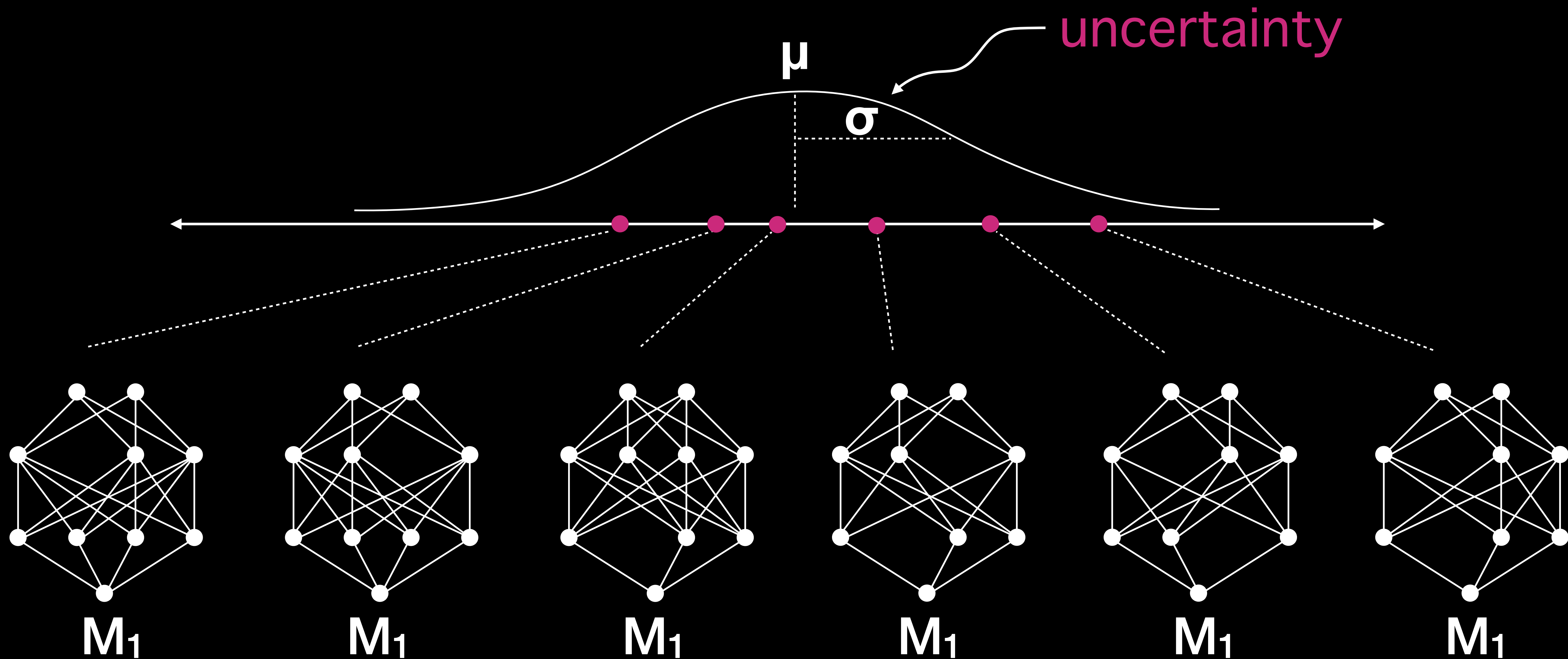


M_1



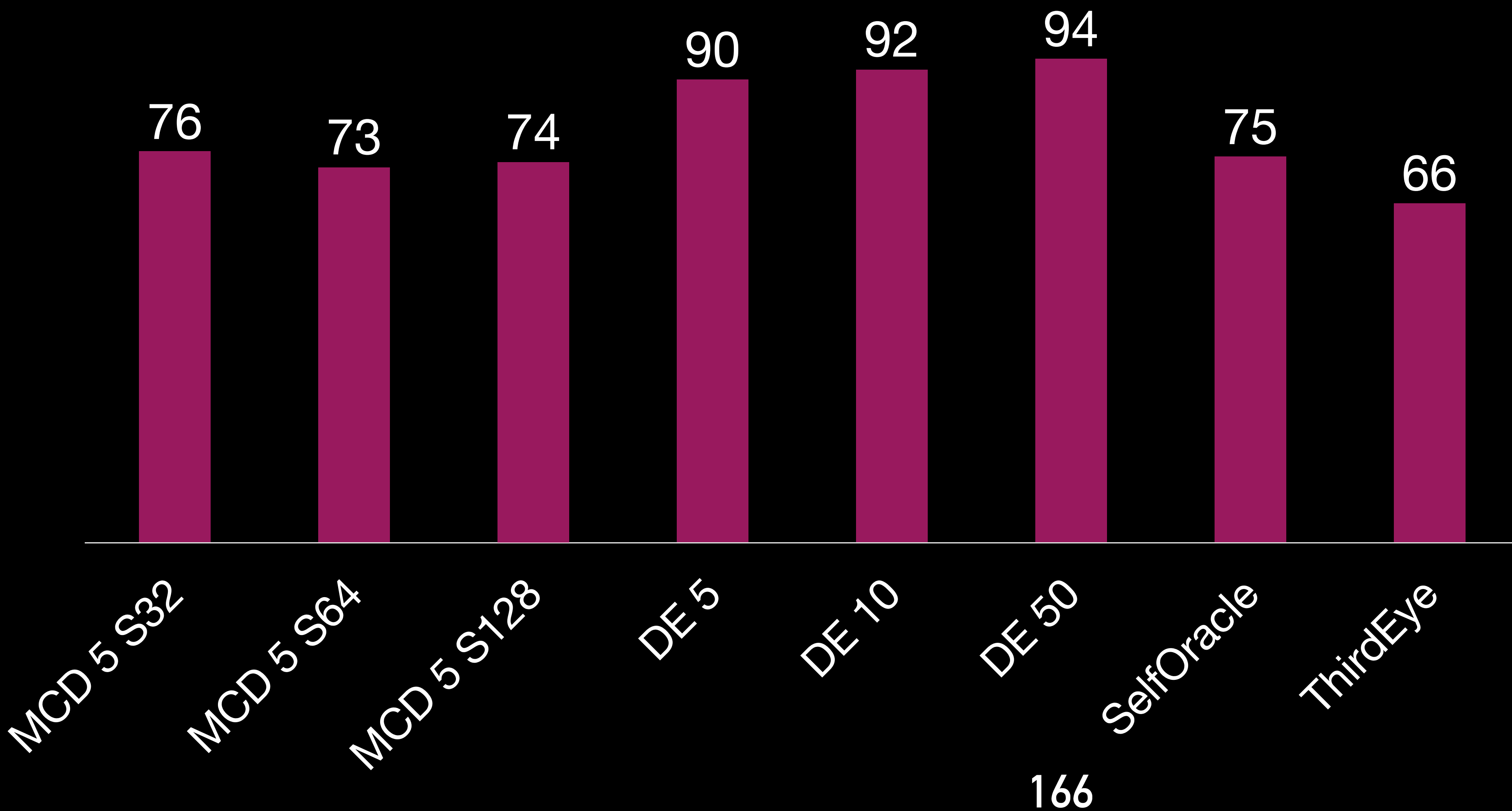
M_1

Monte-Carlo Dropout



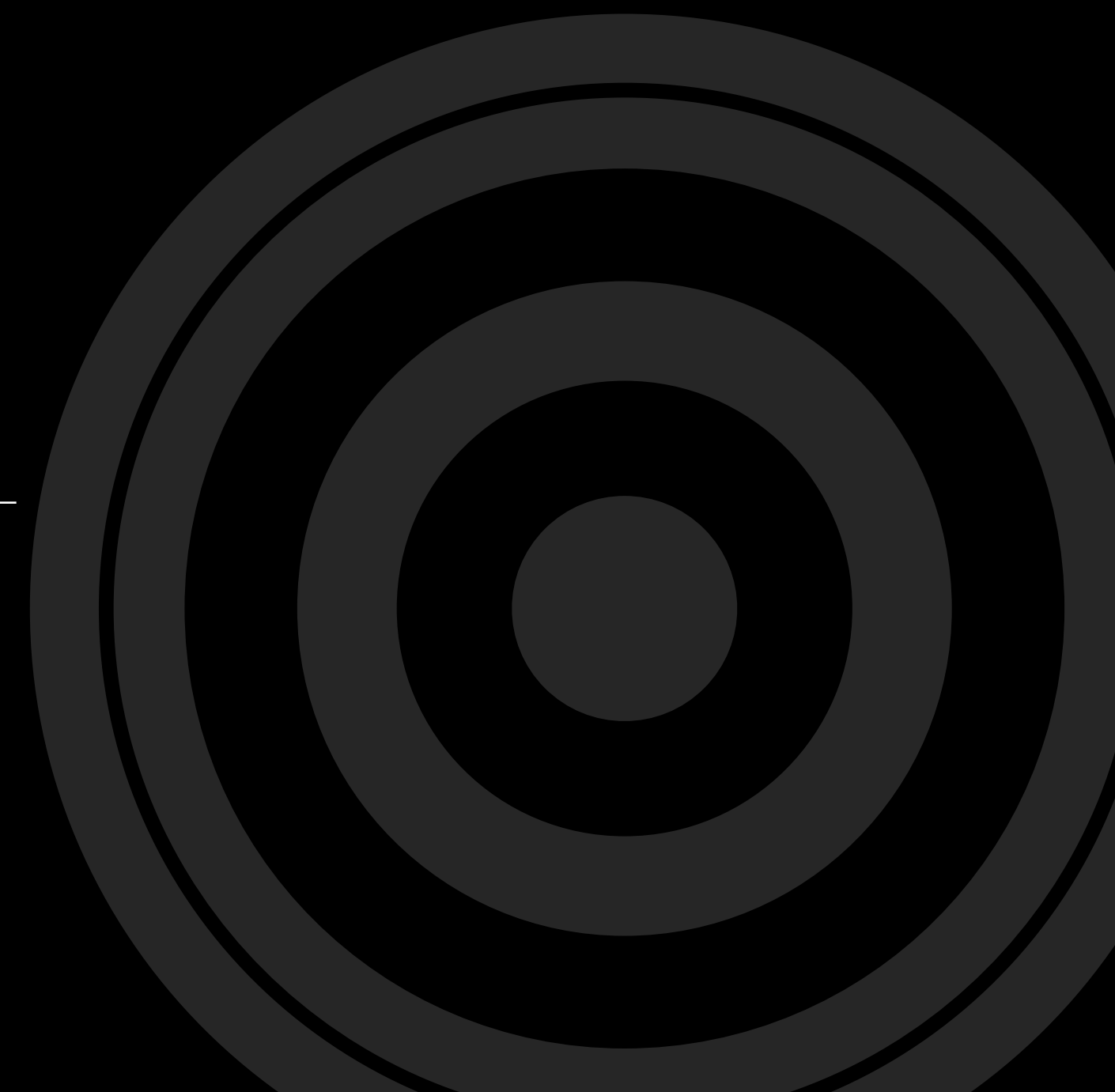
Is our approach **effective**
in **predicting** misbehaviours ?

■ F-3



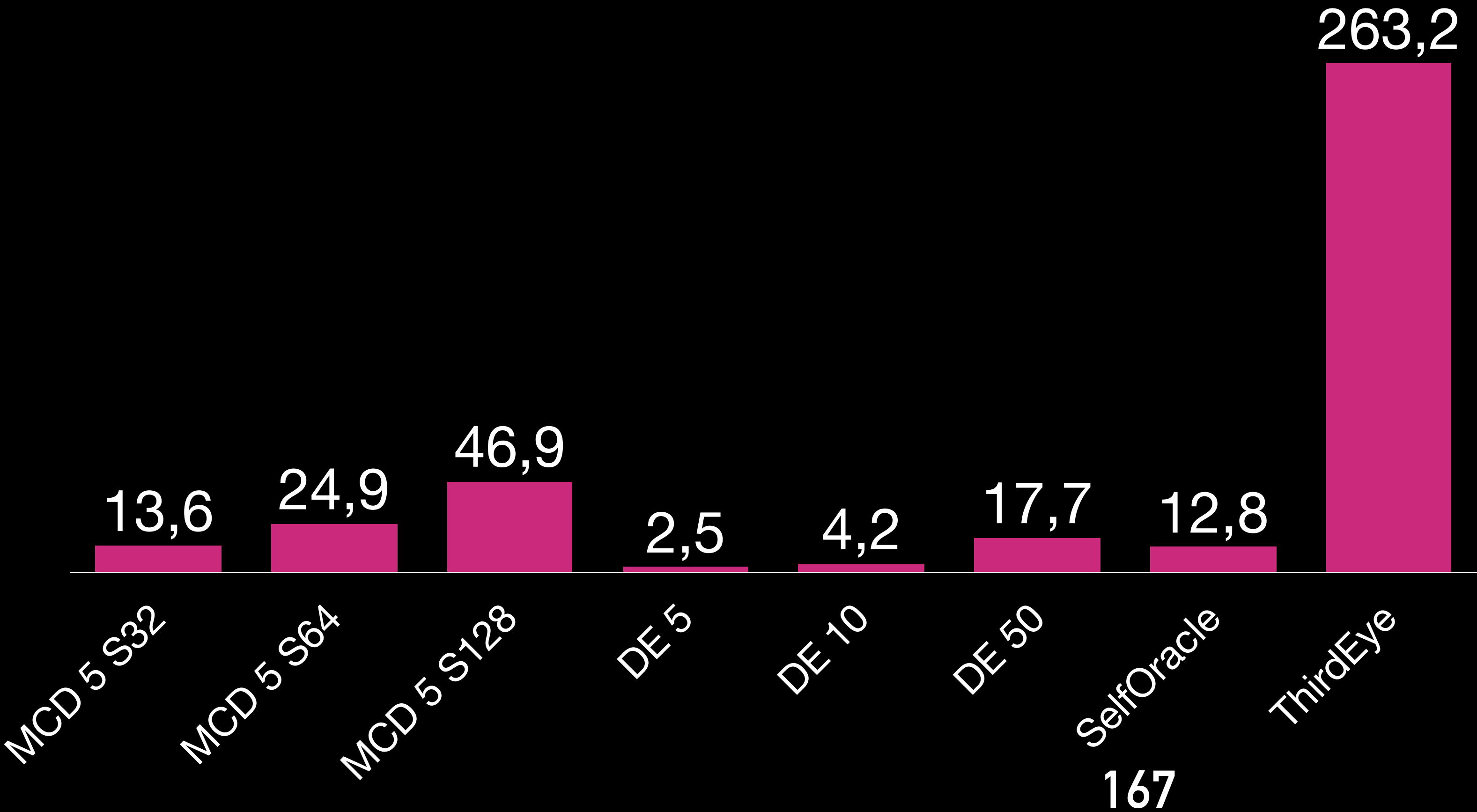
MCD 5 S32:
dropout rate = 5%
S32 = 32 samples

DE 5:
five models



Is our approach efficient
in predicting misbehaviours ?

ms/inference (lower is better)



*Can we **predict**
unexpected conditions*



*Uncertainty Quantification
enables **fast** and **accurate**
online misbehaviour prediction*



How to design them in production?



How to design them in production?



Steps ahead



Testing and Evaluation of Autonomous Driving Systems

From Simulated to Real-world Test Environments

Andrea Stocco

SIESTA summer school

02.09.2026

