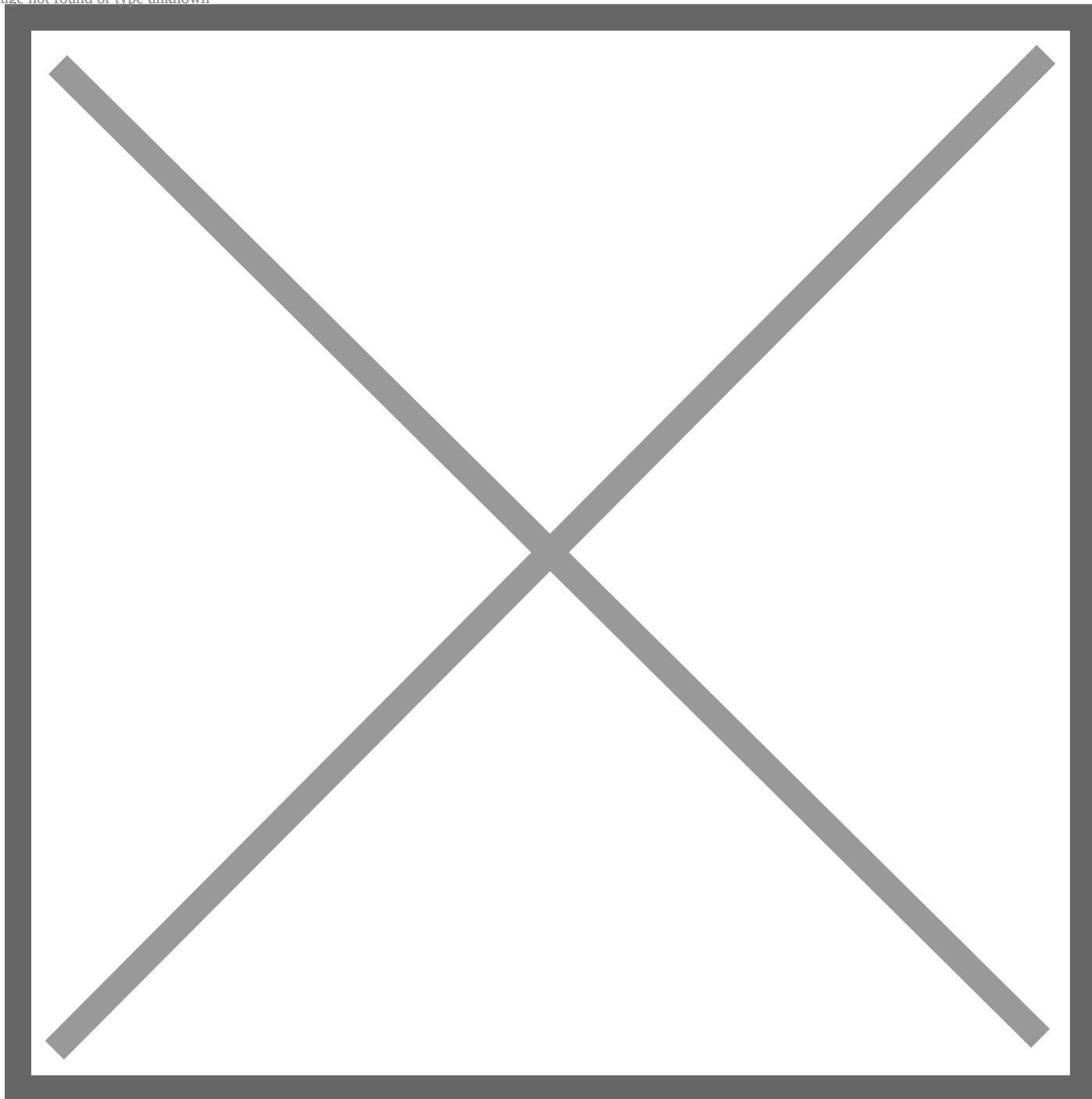


# MIT and Toyota Release Innovative New Visual Open Data to Accelerate Autonomous Driving Research

June 15, 2020

Image not found or type unknown



**CAMBRIDGE, MA (June 15, 2020)** — How can we train self-driving vehicles to have a deeper awareness of the world around them? Can computers learn from past experiences to recognize future patterns that can help them safely navigate new and unpredictable situations?

These are some of the questions researchers from the Massachusetts Institute for Technology (MIT) AgeLab at the MIT Center for Transportation & Logistics and the Toyota Collaborative Safety Research Center (CSRC) are trying to answer by sharing an innovative new open dataset called DriveSeg.

Through the release of DriveSeg, MIT and Toyota are working to advance research in autonomous driving systems that, much like human perception, perceive the driving environment as a continuous flow of visual information.

“In sharing this dataset, we hope to encourage researchers, the industry, and other innovators to develop new insight and direction into temporal AI modeling that enables the next generation of assisted driving and automotive safety technologies,” says Bryan Reimer, Principal Researcher. “Our long-standing working relationship with Toyota CSRC has enabled our research efforts to impact future safety technologies.”

“Predictive power is an important part of human intelligence,” says Rini Sherony, Toyota Collaborative Safety Research Center’s Senior Principal Engineer. “Whenever we drive, we are always tracking the movements of the environment around us to identify potential risks and make safer decisions. By sharing this dataset, we hope to accelerate research into autonomous driving systems and advanced safety features that are more attuned to the complexity of the environment around them.”

To date, self-driving data made available to the research community have primarily consisted of troves of static, single images that can be used to identify and track common objects found in and around the road, such as bicycles, pedestrians or traffic lights through the use of “bounding boxes.” By contrast, DriveSeg contains more precise, pixel-level representations of many of these same common road objects, but through the lens of a continuous video driving scene. This type of full scene segmentation can be particularly helpful for identifying more amorphous objects – such as road construction and vegetation – that do not always have such defined and uniform shapes.

According to Sherony, video-based driving scene perception provides a flow of data that more closely resembles dynamic, real-world driving situations. It also allows researchers to explore data patterns as they play out over time, which could lead to advances in machine learning, scene understanding and behavioral prediction.

DriveSeg is available for free and can be used by researchers and the academic community for non-commercial purposes [here](#). The data is comprised of two parts. DriveSeg (manual) is 2 minutes and 47 seconds of high-resolution video captured during a daytime trip around the busy streets of Cambridge, Massachusetts. The video’s 5,000 frames are densely annotated manually with per-pixel human labels of 12 classes of road objects.

DriveSeg (Semi-auto) is 20,100 video frames (67 × 10 second video clips) drawn from MIT Advanced Vehicle Technologies (AVT) Consortium data. DriveSeg (Semi-auto) is labeled with the same pixel-wise semantic annotation as DriveSeg (manual) except annotations were completed through a novel semiautomatic annotation approach developed by MIT. This approach leverages both manual and computational efforts to coarsely annotate data more efficiently and at lower cost than manual annotation. This data set was created to assess the feasibility of annotating a wide range of real-world driving scenarios and assess the potential of training vehicle perception systems on pixel labels created through AI-based labeling systems.

To learn more about the technical specifications and permitted use?cases for the data, [click here](#).

For more information on Toyota Collaborative Safety Research Center, [click here](#).