

Generating Realistic Traffic Scenarios: A Deep Learning Approach Using Generative Adversarial Networks (GANs)

Md Shadab Alam¹, Marieke H. Martens^{1,2}, and Pavlo Bazilinskyy¹

¹Department of Industrial Design, Eindhoven University of Technology, Eindhoven,
The Netherlands

²TNO, Helmond, The Netherlands

ABSTRACT

Traffic simulations are crucial for testing systems and human behaviour in transportation research. This study investigates the potential efficacy of Unsupervised Recycle Generative Adversarial Networks (Recycle-GANs) in generating realistic traffic videos by transforming daytime scenes into nighttime environments and vice-versa. By leveraging Unsupervised Recycle-GANs, we bridge the gap between data availability during day and night traffic scenarios, enhancing the robustness and applicability of deep learning algorithms for real-world applications. GPT-4V was provided with two sets of six different frames from each day and night time from the generated videos and queried whether the scenes were artificially created based on lightning, shadow behaviour, perspective, scale, texture, detail and presence of edge artefacts. The analysis of GPT-4V output did not reveal evidence of artificial manipulation, which supports the credibility and authenticity of the generated scenes. Furthermore, the generated transition videos were evaluated by 15 participants who rated their realism on a scale of 1 to 10, achieving a mean score of 7.21. Two persons identified the videos as deep-fake generated without pointing out what was fake in the video; they did mention that the traffic was generated.

Keywords: Generative adversarial networks, Future traffic, Deep learning, Traffic modelling, Diurnal traffic behaviour

INTRODUCTION

In contemporary transportation research, data collection efforts encompass diverse traffic scenarios under various conditions, often involving instrumented vehicles equipped with costly sensors such as cameras and LiDAR. Datasets like KITTI Vision Benchmark (Geiger et al., 2012), NuScenes (Caesar et al., 2020), One Thousand and One Hours (Mao et al., 2021), Pedestrian Intention Estimation (PIE) (Rasouli et al., 2019), Waymo Open Dataset (Sun et al., 2020), ApolloScape Auto (Huang et al., 2020) and Cityscapes (Cordts et al., 2016) are benchmarks for numerous computer vision and automated driving-related tasks. These datasets have been used in various studies to get specific insights. Kooijman (2021) used PIE in a crowdsourced experiment, exploring the impact of objective in-scene

features on driver perceptions during interactions with pedestrians. Similarly, Gu et al. (2020) trained an LSTM-based automated driving model based on the Waymo Open Dataset to imitate the behaviour of Waymo's self-driving model.

Many datasets exhibit a bias toward data collected under sunny daylight conditions while lacking sufficient representation of nighttime or other seasonal scenarios. For example, the KITTI dataset comprises annotations of more than 200,000 3D objects captured in cluttered urban scenarios but have only scenes from sunny daylight environments. Similarly, dataset One Thousand and One Hours (Mao et al., 2021) comprises 1,118 hours of footage collected in Palo Alto, CA, USA, which were collected with 20 automated vehicles driving along a fixed route. However, it also lacks adequate representation of nighttime or diverse weather conditions. The Cityscapes dataset (Cordts et al., 2016), which includes 5,000 finely annotated urban stereo video sequences collected across 50 cities, primarily in Germany but also in neighbouring countries, suffers from similar limitations. The datasets AppoloScape (Huang et al., 2020) and KAIST (Hwang et al., 2015) do not have scenes from night-time, although they have scenes from rainy weather conditions.

Both academic and industrial projects often rely on simulations generated by software platforms like Unity (<https://unity.com>) or Unreal Engine (<https://unrealengine.com>) or use footage from public sources on the Internet to conduct experiments. For instance, Bazilinskyy et al. (2022) presented 1,438 participants in a crowdsourcing study with scenes developed from Unity. The participants viewed 227 textual external Human Machine Interface (eHMI) on a vehicle and were asked to rate the stimuli based on their confidence to follow the prompt displayed on the eHMI. Tran et al. (2024) used Unity to develop scenes for understanding the influence of one pedestrian over another in the presence of a vehicle equipped with eHMI. Similarly, Alam et al. (2024) used scenes developed in Unreal Engine to understand user preference to board a shuttle bus. Rasouli et al. (2017) curated a dataset from prompting participants to predict pedestrian intentions when crossing the road. Companies like BMW (Group, 2017) run their cars for several million kilometres in simulation for testing the vehicle.

Recently researchers have started to look forward toward deep learning architectures such as GANs (Goodfellow et al., 2020) and diffusion models (Song and Ermon, 2019) for transforming scenes. Studies by Parmar et al. (2024), Zhu et al. (2017), Alam et al. (2024), Isola et al. (2017) and Anoosheh et al. (2019) showed the transformation of scenes from one environment to another. For example, Anoosheh et al. (2019) proposed ToDayGAN a GAN model which can translate nighttime driving images to daytime representation allowing robots to determine their position and orientation in the world. Still, this work is limited to transforming from scenario from night to day only. Similarly, Parmar et al. (2024) translated fumes from day to night and vice versa and introduced clear to rainy weather transformation. However, their transformations suffered from a lack of tonal constraints due

to conflicts between the noise map and the input conditioning image in one-step models, which impacted the tonal coherence and structural accuracy of the output.

Aim of the Study

This study investigates the effectiveness of Unsupervised Recycle-GANs (Wang et al., 2022) in transforming traffic scenes between daytime and nighttime environments. The primary objective is to evaluate the model's ability to generate realistic traffic scenarios under varying lighting conditions, ensuring the visual quality and fidelity of the generated outputs. To achieve this, we will first collect a training dataset from a publicly available street camera stream on YouTube (<https://www.youtube.com>), comprising traffic scenes captured during both day and night. Once the Unsupervised Recycle-GAN is trained, the model will be tested using new video footage from the same camera source to validate its generalisation capabilities. The generated traffic scenes will undergo a two-step evaluation process. First, participants will provide subjective feedback on the perceived realism of the day-to-night and night-to-day transformations. This human evaluation will focus on how convincingly the transitions emulate real-world traffic conditions. Second, an evaluation will be performed using GPT-4V to provide a structured and objective assessment of the generated scenes.

METHOD

The research was approved by the Human Research Ethics Committee of the Eindhoven University of Technology. We employed a dataset from a live camera stream on YouTube (<https://www.youtube.com/watch?v=JbnJAsk1zII>) captured on Gangnam Street in Seoul, South Korea. The footage included one hour of daytime and one hour of nighttime scenes, recorded on 5 April 2024, 16:00–17:00 (GMT + 9) and 5 April 2024, 20:00–21:00 (GMT + 9), respectively. We strategically selected these times to maximise pedestrian activity, as the late afternoon typically sees increased foot traffic, while choosing too late in the evening might result in fewer people on the streets, potentially impacting the richness and diversity of the dataset. The videos are available in the supplementary material. To ensure thorough coverage of both daylight and nighttime conditions, we divided the footage into training and validation datasets, dedicating 80% to training and 20% to validation. This dataset served as the foundation for the training and evaluation of our proposed models, enabling the exploration of traffic behaviour under varying lighting conditions and environmental settings. We again recorded footage on 13 April 2024 for testing purposes. This time, the recordings were limited to 10 minutes each for both daytime and nighttime, with the time slots selected randomly to minimise any potential bias in the dataset.

We used Unsupervised Recycle-GANs architecture (Wang et al., 2022) to train the network. The contemporary difference between Recycle-GANs and unsupervised Recycle-GANs is that they incorporate tonal constraints in the learning process, specifically focusing on enhancing the visual quality and

the realism of the generated images. This distinction is crucial for generating traffic scenes with high fidelity, as it ensures that the synthetic images closely resemble the characteristics of real-world traffic scenarios. Additionally, the utilisation of Recycle-GANs facilitates the preservation of essential features such as vehicle shapes, colours, and movement patterns during the generation process, thus contributing to the overall effectiveness of the framework in simulating realistic traffic dynamics.

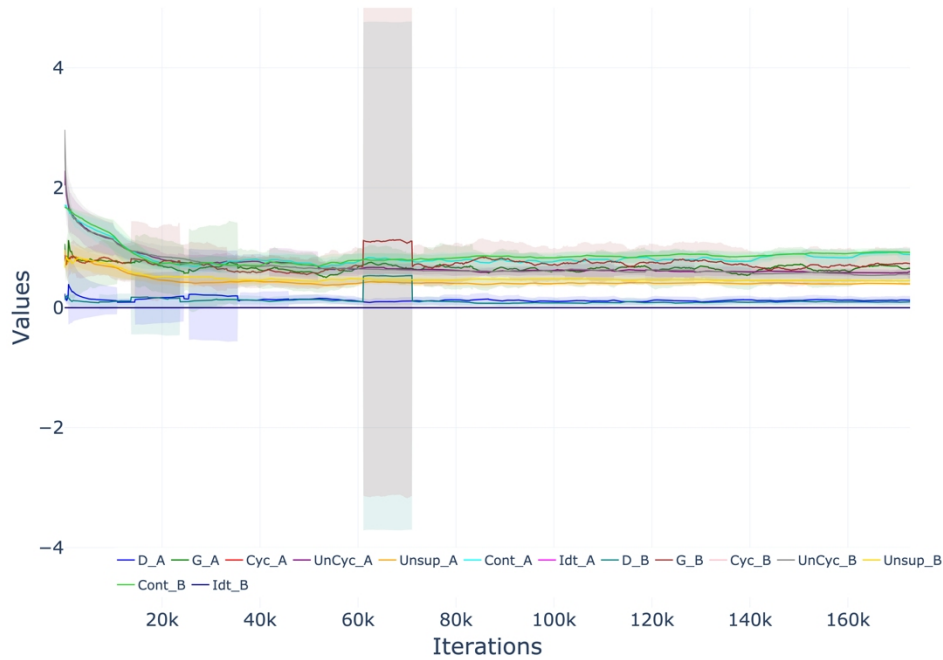


Figure 1: Training loss for the model. D_A and D_B are discriminator losses, while G_A and G_B are generator losses for domains A and B. Cyc_A and Cyc_B enforce cycle consistency, ensuring images translated between domains and back remain unchanged. UnCyc_A, UnCyc_B, Unsup_A, and Unsup_B promote temporal and spatial consistency, respectively. Cont_A and Cont_B preserve semantic content, and Idt_A and Idt_B maintain identity consistency within each domain.

Upon obtaining the results, the subsequent step involved assessing the veracity of the images. To achieve this, we utilised GPT-4V. This has been done in numerous research projects (De Winter, 2024; Driessen et al., 2024; Tabone and De Winter, 2023; Talan and Kalinkara, 2023; Wardat et al., 2023). For example, Driessen et al. (2024) assessed GPT-4V's ability to predict human-perceived risk levels in traffic images, utilising 210 static images rated by approximately 650 individuals. They found that repeating prompts under identical conditions, varying prompt text, and incorporating object detection features alongside GPT-4V-based risk ratings significantly enhance model validity. Typically, genuine photographs captured by cameras exhibit consistent lighting, shadows that align with light sources, accurate perspectives and scales, and natural textures. Deviations from these elements within an image usually indicate that it was edited or completely synthesised.

While these indicators are not solely characteristic of fabricated images, they represent anomalies that are generally absent in authentic, unaltered photographs of real-world scenes. In authentic images, all elements are expected to align cohesively with the physics of light and space. During our assessment, we searched for discrepancies from this natural coherence to determine if an image appeared fabricated. We inputted 12 distinct scenes into GPT-4V and requested its assessment regarding the authenticity of the images. To ensure a comprehensive evaluation and to check the consistency of the response from GPT-4V's we divided the 12 scenes into 2 batches comprising 6 scenes, each containing 3 day and 3 night scenarios. In both cases, GPT-4V was prompted with *“Based on the following criteria, could you determine if these images are artificially created? 1. Uniformity of lighting 2. Shadow behavior 3. Perspective and scale 4. Texture and detail 5. Presence of edge artefacts”*.

Then, secondly, a total of 15 participants (7 males and 8 females) took part in the questionnaire. The participants had a mean age of 24.73 years ($SD = 5.28$). The nationalities of the participants were Dutch (11), Chinese (2), Colombian (1), and Indonesian (1). Each participant was shown two test videos: first depicting a transformation from night to day and the other from day to night. The length of both videos was kept consistent at 1 minute. To ensure unbiased responses, participants were not informed beforehand whether the videos were real or artificially generated. After viewing both videos, participants were asked to respond to the following questions: (Q1) *“What is your overall view on the video?”* and (Q2) *“Was there anything in the video that seemed unusual or didn't feel quite right?”*. The participants were also asked to rate the overall quality of the video (lighting, movement, details) on a scale of 1 to 10, with 10 representing excellent quality. Once the participants provided their initial evaluations and ratings, they were informed that the videos were artificially created. Following this, a feedback session was conducted to gather further insights. Participants were asked: (Q3) *“Now that you know the video was artificial, did anything stand out to you that seemed noticeably fake?”* and (Q4) *“If you hadn't been told the video was artificial, do you think you would have noticed?”*

RESULTS

After completing the training process, the GANs neural network underwent testing using footage captured on 13 April 2024. The training loss averaged over 100 iterations and the variance are shown in Figure 1. The resulting videos shows 10-minute-long daytime and nighttime scenarios. The videos are available in the supplementary material. To visualise the efficacy of the trained model, Figure 2 presents a single-frame comparison between the daytime and nighttime conditions alongside its transformation through the GANs. Notably, the GANs adeptly transpose scenes from one lighting condition to another, as evidenced by the transition in the transformed images.

The generated videos were evaluated with GPT-4V. Table 1 shows the received feedback. For the first set of six images, GPT-4V indicated: *“Based*

on these observations, there are no definitive indications that these images are artificially created”. Interestingly, GPT-4V identified each image and gave comments on daylight images and nighttime. For the second set of six images, GPT-4V commented: “Overall, these images also do not exhibit clear signs of artificial creation upon visual inspection. They appear to maintain consistent lighting, shadow behaviour, perspective, and detail that one would expect from unaltered photos”. The detailed response is shown in Table 1.

The generated videos were also presented to the participants for evaluation. The results of this evaluation are available in the supplementary material. When asked for their overall view of the videos (Q1), all 15 participants focused on the chaotic nature of the traffic depicted in the scenes. All participants were unfamiliar with Korean streets and 8 of them commented on the perceived lack of safety for pedestrians due to the seemingly random paths followed by vehicles and pedestrians. A notable observation among participants was the attention drawn to the zebra crossings and road markings, particularly the yellow and white colours.

When participants were asked whether anything in the video seemed unusual or did not feel quite right (Q2), their responses primarily continued to emphasise the chaotic traffic. However, one participant pointed out the presence of halos around some people in the video as a potentially unusual element.

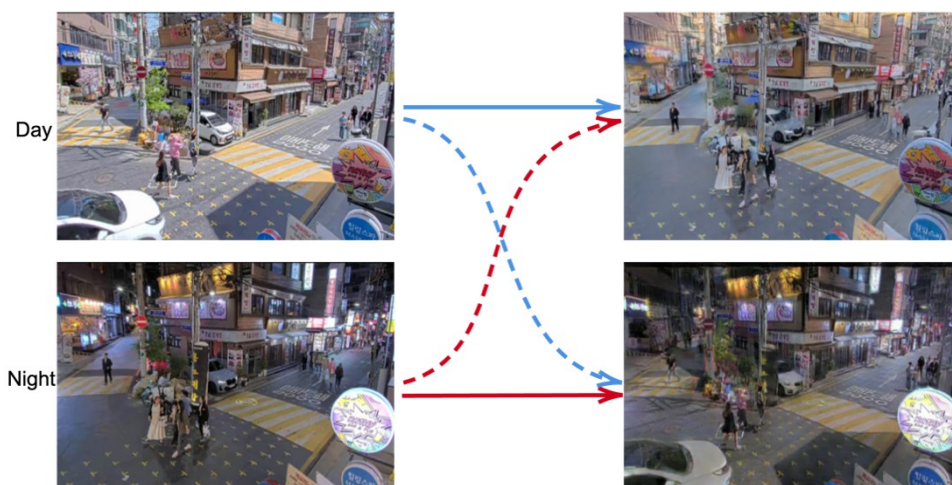


Figure 2: Dynamic scene translation with Unsupervised Recycle-GANs: This frame from a dynamic video exemplifies the application of Unsupervised Recycle-GANs in seamlessly transforming day scenes into night and vice versa. The images on the left depict the original footage, while those on the right are generated via Unsupervised Recycle-GANs.

The participants were then asked to rate the overall quality of the video on a scale of 1 to 10, with 10 representing excellent quality. The average score was 7.21 (SD = 0.81). After this initial evaluation, participants were informed that the videos were generated using GANs. This revelation (Q3) prompted varied reactions: two participants expressed their surprise out loud.

5 participants mentioned they had seen some inconsistencies related to the lighting conditions, but 7 participants continued to focus on the traffic and zebra crossings in their feedback.

Table 1. GPT-4V evaluation of the 12 images for signs of artificial creation, examining criteria such as uniformity of lighting, shadow behaviour, perspective and scale, texture and detail, and presence of edge artefacts.

Queries	Response 1	Response 2
Uniformity of Lighting	- In the night images, the lighting does appear to have a consistent source, with illuminated signs and streetlights casting appropriate brightness on nearby objects. - In the day images, the sunlight appears natural, casting shadows that match the direction of light typically seen during daytime.	- There is consistency in the lighting, especially visible in the day images with the sunlight. - The nighttime images have varied lighting sources which seem to reflect naturally on the surfaces around
Shadow Behavior	- Shadows in the daytime images are consistent with sunlight direction, and the shadows cast by people and objects in the street appear correct. - In the nighttime images, shadows are less visible but where they are seen, they seem to correspond to the artificial light sources.	- Shadows cast by pedestrians and objects in the day images align with the sun's position. - At night, the shadows that are visible seem to be cast by the streetlights appropriately, although it's more challenging to discern due to lower light levels.
Perspective and Scale	- Buildings, people, and street elements are in perspective, with no apparent distortion that would suggest artificial manipulation. - The scale of objects seems consistent; larger objects are closer, smaller objects are further away, suggesting a correct perspective.	- Perspective and scale look consistent. The buildings and street elements decrease in size appropriately as they recede into the distance. - People and objects closer to the camera are larger, as would be expected in a three-dimensional space.
Texture and Detail	- There is a noticeable granularity to the images, especially visible in the road surface and building facades, which suggests that texture is well represented. - Details such as street signs, window patterns, and the texture of clothing on people appear natural and varied without obvious repetition.	- The textures, such as those on the road surface and the building facades, are clear and varied. - Details are consistent across different elements within each image, suggesting a natural scene.
Presence of Edge Artifacts	No obvious edge artifacts, such as halos or unnaturally sharp edges, are visible at this resolution, indicating the images might be authentic or well edited.	No clear edge artifacts are visible upon this inspection. However, without zooming in and examining at a higher resolution, it's possible to miss subtle signs of digital manipulation.

Finally, participants were asked whether they would have noticed that the videos were artificially created if they had not been told (Q4). Only two participants claimed that they would have identified the artificial nature of the videos without being informed although they both gave a 7.5 rating which is above average, while the rest indicated that they would not have noticed.

DISCUSSION

In this study, we demonstrated the effectiveness of Unsupervised Recycle-GANs for traffic scene transformation across different times of the day. Our approach addresses the challenge of the lack of datasets for training machine learning algorithms for transportation research, particularly in low-light conditions such as nighttime scenes. By leveraging Recycle-GANs, we bridge the gap between data availability during day and night scenarios, enhancing the robustness and applicability of traffic analysis models.

The findings demonstrate the potential of Unsupervised Recycle-GANs in generating realistic traffic videos with minimal perceptible flaws. The analysis conducted using GPT-4V corroborated the authenticity of the generated scenes, finding no evidence of artificial manipulation based on factors such as lighting, shadow behaviour, and texture details. These results were further supported by human evaluation, where most of the 15 participants were unable to discern the artificial nature of the videos without prior disclosure, highlighting the model's effectiveness in creating convincing day-to-night and night-to-day transitions. While some participants noted specific visual elements, such as zebra crossings or the chaotic nature of traffic, only a few identified subtle artefacts like halos around pedestrians or lighting inconsistencies. Notably, two participants correctly identified the videos as artificially generated but still rated them highly (7.5 each), despite initially attributing the artificiality to elements such as pedestrians or cars, which were incorrectly perceived as fake. These findings suggest that Unsupervised Recycle-GANs offer a robust approach to bridging the data gap between day and night traffic scenarios, enhancing the applicability of deep learning in realistic traffic simulations.

In future research, this study can be extended to generate custom videos featuring diverse scenarios with varying densities of cars and pedestrians. Additionally, there is scope for integrating cross-cultural perspectives into the training process, encompassing traffic conditions from different regions worldwide. Furthermore, a crucial aspect of enhancement lies in incorporating background sounds, such as traffic honks and ambient noise, into the generated videos. This integration would enhance the realism and immersion of the simulated traffic scenes, offering a more comprehensive dataset for analysis and training of machine learning algorithms. Additionally, exploring techniques for fine-tuning the generated videos to specific cultural and geographical contexts can further enhance the utility and accuracy of the generated traffic simulations. Comparing our approach with OpenAI's SORA (<https://openai.com/sora>) would provide valuable insights and contribute to advancing traffic simulation technology.

SUPPLEMENTARY MATERIAL

Videos used to test, train and validate the network, source code in Python, and anonymised results of the evaluation with participants can be found at <https://doi.org/10.4121/80c664cb-a4b5-4eb1-bc1c-666349b1b927>. A maintained version of the code is available at <https://github.com/Shaadalam9/gans-traffic>.

REFERENCES

- Alam, M. S., Parmar, S. H., Martens, M. H., Bazilinskyy, P., 2025. Deep Learning Approach for Realistic Traffic Video Changes Across Lighting and Weather Conditions. Presented at the 8th International Conference on Information and Computer Technologies (ICICT), Hilo, Hawaii, USA.
- Anoosheh, A., Sattler, T., Timofte, R., Pollefeys, M., Gool, L. V., 2019. Night-to-day image translation for retrieval-based localization, in: 2019 International Conference on Robotics and Automation (ICRA). IEEE Press, Montreal, QC, Canada, pp. 5958–5964. <https://doi.org/10.1109/ICRA.2019.8794387>
- Bazilinskyy, P., Dodou, D., De Winter, J. C. F., 2022. Crowdsourced assessment of 227 text-based eHMI for a crossing scenario, in: Proceedings of International Conference on Applied Human Factors and Ergonomics (AHFE). New York, USA. <https://doi.org/10.54941/ahfe1002444>
- Caesar, H., Bankiti, V., Lang, A. H., Vora, S., Liong, V. E., Xu, Q., Krishnan, A., Pan, Y., Baldan, G., Beijbom, O., 2020. nuScenes: A multimodal dataset for autonomous driving, in: 2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 11618–11628. <https://doi.org/10.1109/CVPR.42600.2020.01164>
- Cordts, M., Omran, M., Ramos, S., Rehfeld, T., Enzweiler, M., Benenson, R., Franke, U., Roth, S., Schiele, B., 2016. The cityscapes dataset for semantic urban scene understanding, in: 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 3213–3223. <https://doi.org/10.1109/CVPR.2016.350>
- De Winter, J. C. F., 2024a. Can ChatGPT pass high school exams on English language comprehension? *Int. J. Artif. Intell. Educ.* 34, 915–930. <https://doi.org/10.1007/s40593-023-00372-z>
- De Winter, J. C. F., 2024b. Can ChatGPT be used to predict citation counts, readership, and social media interaction? An exploration among 2222 scientific abstracts. *Scientometrics* 129, 2469–2487. <https://doi.org/10.1007/s11192-024-04939-y>
- De Winter, J. C. F., Hoogmoed, J., Stapel, J., Dodou, D., Bazilinskyy, P., 2023. Predicting perceived risk of traffic scenes using computer vision. *Transp. Res. Part F Traffic Psychol. Behav.* 93, 235–247. <https://doi.org/10.1016/j.trf.2023.01.014>
- Driessen, T., Dodou, D., Bazilinskyy, P., Winter, J. de, 2024. Putting ChatGPT vision (GPT-4V) to the test: risk perception in traffic images. *R. Soc. Open Sci.* <https://doi.org/10.1098/rsos.231676>
- Geiger, A., Lenz, P., Urtasun, R., 2012. Are we ready for autonomous driving? The KITTI vision benchmark suite, in: 2012 IEEE Conference on Computer Vision and Pattern Recognition. Presented at the 2012 IEEE Conference on Computer Vision and Pattern Recognition, pp. 3354–3361. <https://doi.org/10.1109/CVPR.2012.6248074>
- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., Bengio, Y., 2020. Generative Adversarial Networks. *Commun ACM* 63, 139–144. <https://doi.org/10.1145/3422622>

- Group, B., 2017. In Sprints Towards Autonomous Driving.
- Gu, Z., Li, Z., Di, X., Shi, R., 2020. An LSTM-Based Autonomous Driving Model Using a Waymo Open Dataset. *Appl. Sci.* 10, 2046. <https://doi.org/10.3390/ap10062046>
- Huang, X., Wang, P., Cheng, X., Zhou, D., Geng, Q., Yang, R., 2020. The ApolloScape open dataset for autonomous driving and its application. *IEEE Trans. Pattern Anal. Mach. Intell.* 42, 2702–2719. <https://doi.org/10.1109/TPAMI.2019.2926463>
- Hwang, S., Park, J., Kim, N., Choi, Y., Kweon, I. S., 2015. Multispectral pedestrian detection: Benchmark dataset and baseline, in: 2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 1037–1045. <https://doi.org/10.1109/CVPR.2015.7298706>
- Isola, P., Zhu, J.-Y., Zhou, T., Efros, A. A., 2017. Image-to-image translation with conditional adversarial networks. Presented at the 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), IEEE Computer Society, pp. 5967–5976. <https://doi.org/10.1109/CVPR.2017.632>
- Kooijman, B., 2021. The identification of factors affecting drivers' perceived risk in pedestrian-vehicle interaction: A crowdsourcing study.
- Mao, J., Niu, M., Jiang, C., Liang, H., Chen, J., Liang, X., Li, Y., Ye, C., Zhang, W., Li, Z., Yu, J., Xu, H., Xu, C., 2021. One Million Scenes for Autonomous Driving: ONCE Dataset. <https://doi.org/10.48550/arXiv.2106.11037>
- Parmar, G., Park, T., Narasimhan, S., Zhu, J.-Y., 2024. One-step image translation with text-to-image models. <https://doi.org/10.48550/arXiv.2403.12036>
- Rasouli, A., Kotseruba, I., Kunic, T., Tsotsos, J., 2019. PIE: A large-scale dataset and models for pedestrian intention estimation and trajectory prediction, in: 2019 IEEE/CVF International Conference on Computer Vision (ICCV). pp. 6261–6270. <https://doi.org/10.1109/ICCV.2019.00636>
- Song, Y., Ermon, S., 2019. Generative modeling by estimating gradients of the data distribution, in: Proceedings of the 33rd International Conference on Neural Information Processing Systems. Curran Associates Inc., Red Hook, NY, USA, pp. 11918–11930. <https://doi.org/10.48550/arXiv.1907.05600v3>
- Sun, P., Kretschmar, H., Dotiwalla, X., Chouard, A., Patnaik, V., Tsui, P., Guo, J., Zhou, Y., Chai, Y., Caine, B., others, 2020. Scalability in perception for autonomous driving: Waymo open dataset, in: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 2446–2454. <https://doi.org/10.48550/arXiv.1912.04838>
- Tabone, W., De Winter, J. C. F., 2023. Using ChatGPT for human–computer interaction research: a primer. *R. Soc. Open Sci.* <https://doi.org/10.1098/rsos.231053>
- Talan, T., Kalinkara, Y., 2023. The Role of Artificial Intelligence in Higher Education: ChatGPT Assessment for Anatomy Course. *Int. J. Manag. Inf. Syst. Comput. Sci.* 7, 33–40. <https://doi.org/10.33461/uybisbbd.1244777>
- Wang, K., Akash, K., Misu, T., 2022. Learning Temporally and Semantically Consistent Unpaired Video-to-video Translation Through Pseudo-Supervision From Synthetic Optical Flow. <https://doi.org/10.48550/arXiv.2201.05723>
- Wardat, Y., Tashtoush, M. A., AlAli, R., Jarrah, A. M., 2023. ChatGPT: A revolutionary tool for teaching and learning mathematics. *Eurasia J. Math. Sci. Technol. Educ.* 19, em2286. <https://doi.org/10.29333/ejmste/13272>
- Zhu, J.-Y., Park, T., Isola, P., Efros, A. A., 2017. Unpaired image-to-image translation using cycle-consistent adversarial networks, in: Proceedings of the IEEE International Conference on Computer Vision. pp. 2223–2232. <https://doi.org/10.48550/arXiv.1703.10593>