



(51) International Patent Classification:

G06K 9/00 (2006.01)

G06K 9/62 (2006.01)

(21) International Application Number:

PCT/EP2019/055786

(22) International Filing Date:

07 March 2019 (07.03.2019)

(25) Filing Language:

English

(26) Publication Language:

English

(71) Applicants: **TOYOTA MOTOR EUROPE** [BE/BE]; Avenue du Bourget 60, 1140 BRUSSELS (BE). **ETH ZURICH** [CH/CH]; THE SWISS FEDERAL INSTITUTE OF, TECHNOLOGY ZURICH, Raemistrasse 101, 8092 ZURICH (CH).

(72) Inventors: **VIGNARD, Nicolas**; c/o Toyota Motor Europe, European Technical Administration, Division, Avenue du Bourget 60, 1140 BRUSSELS (BE). **DAI, Dengxin**; c/o THE SWISS FEDERAL INSTITUTE OF, TECHNOLOGY ZURICH, Raemistrasse 101, 8092 ZURICH (CH).

HECKER, Simon; c/o THE SWISS FEDERAL INSTITUTE OF, TECHNOLOGY ZURICH, Raemistrasse 101, 8092 ZURICH (CH). **VAN GOOL, Luc**; c/o THE SWISS FEDERAL INSTITUTE OF, TECHNOLOGY ZURICH, Raemistrasse 101, 8092 ZURICH (CH).

(74) Agent: **HEWEL, Christoph** et al.; CABINET BEAU DE LOMENIE, 158 Rue de l'Université, 75340 PARIS CEDEX 07 (FR).

(81) Designated States (unless otherwise indicated, for every kind of national protection available): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JO, JP, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL, SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, ZA, ZM, ZW.

(54) Title: SYSTEM AND METHOD FOR TRAINING A MODEL PERFORMING HUMAN-LIKE DRIVING

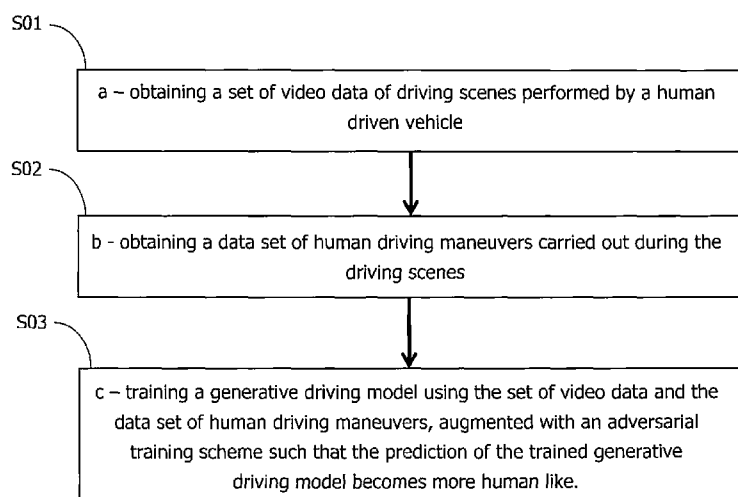


Fig. 1

(57) Abstract: The invention relates to a method and system for training a human-like generative driving model for a vehicle, comprising: a - obtaining (S01) a set of video data of driving scenes performed by a human driven vehicle, b - obtaining (S02) a data set of human driving maneuvers carried out during the driving scenes, c - training (S03) a generative driving model using the set of video data and the data set of human driving maneuvers, wherein the training step (S03) of c is augmented with an adversarial training scheme such that the prediction of the trained generative driving model becomes more human like.

(84) **Designated States** (*unless otherwise indicated, for every kind of regional protection available*): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), Eurasian (AM, AZ, BY, KG, KZ, RU, TJ, TM), European (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

Published:

— *with international search report (Art. 21(3))*

SYSTEM AND METHOD FOR TRAINING A MODEL PERFORMING HUMAN-LIKE DRIVING

FIELD OF THE DISCLOSURE

5 [0001] The present disclosure is related to the field of image processing, in particular to a method for training a human-like generative driving model for a vehicle.

BACKGROUND OF THE DISCLOSURE

10 [0002] The prospect of deploying autonomously driven cars is imminent owing to the advances in perception, robotics and sensor technologies. However, it is believed that autonomous vehicles are more likely to be accepted if they drive accurately, comfortably and drive the same way as human drivers would do. This is especially true for the near future when autonomous vehicles and human-driven vehicles need to share the same road.

15 [0003] Classical approaches require the recognition of all driving-relevant objects, such as lanes, traffic signs, traffic lights, cars and pedestrians, and then perform motion planning, which is further used for final vehicle control, cf. e.g.:

20 C. Urmson, et.al. Autonomous driving in urban environments: Boss and the urban challenge. Journal of Field Robotics Special Issue on the 2007 DARPA Urban Challenge, Part I, 25(8):425–466, June 2008.

25 [0004] These type of systems are sophisticated, represent the current state-of-the-art for autonomous driving, but they are hard to maintain and prone to error accumulation over the pipeline. Most systems also need to use diverse sensors, such as cameras, laser scanners, radar, GPS and high-definition maps.

 [0005] End-to-end mapping methods on the other hand construct a direct mapping from the sensory input to the maneuvers. In this regard the last years have seen tremendous progress in academia on learning driving models, cf. e.g.:

30 F. Codevilla, M. Müller, A. Lopez, V. Koltun, and A. Dosovitskiy. End-to-end driving via conditional imitation learning. 2018, and
 S. Hecker, D. Dai, and L. Van Gool. End-to-end learning of driving models with surround-view cameras and route planners. In ECCV, 2018.

35 [0006] However, many of these systems are deficient in terms of the sensors used, when compared to the driving systems developed by large

companies. For instance, many algorithms only use a front-facing camera. Maps are exploited only for simple directional commands or rendered videos. While these setups are sufficient to allow the community to study many challenges, developing algorithms for fully autonomous cars requires the use of numerical maps of high fidelity.

[0007] Current driving algorithms, e.g. those cited above, mostly treat driving as a regression problem with i.i.d individual training samples, e.g. regressing the low-level steering angle and speed for a given data sample. Yet, driving is a continuous sequence of events over time. Longitudinal and lateral control need to be coupled and these coupled operations need to be combined over time for a comfortable ride. Thus, driving models need to be learned with continuous data sequences and proper passenger comfort measures need to be embedded into the learning system.

[0008] Other contributions have chosen the middle ground between traditional pipe-lined methods and the monolithic end-to-end approach. They learn driving models from compact intermediate representations called affordance indicators such as distance to the front car and existence of traffic light, cf. e.g.

A. Sauer, N. Savinov, and A. Geiger. Conditional affordance learning for driving in urban environments. In Conference on Robot Learning, 2018.

[0009] While research on passenger comfort started to receive some attention, it hardly did so in learning driving models, cf. e.g.:

M. Elbanhawi, M. Simic, and R. Jazar. In the passenger seat: Investigating ride comfort measures in autonomous cars. IEEE Intelligent Transportation Systems Magazine, 7(3):4–17, 2015).

[0010] A large body of work has studied human driving styles, cf. e.g.:

G. A. M. Meiring and H. C. Myburgh. A review of intelligent driving style analysis systems and related artificial intelligence algorithms. In Sensors, 2015.

[0011] Statistical approaches have been employed to evaluate human drivers and to suggest improvements, cf. e.g.:

H. Zhao, H. Zhou, C. Chen, and J. Chen. Join driving: A smart phone-based driving behavior evaluation system. In IEEE Global Communications Conference (GLOBECOM), 2013.

However, human-like driving is hard to quantify.

SUMMARY OF THE DISCLOSURE

[0012] Currently, it remains desirable to provide a system and a method for training a human-like generative driving model for a vehicle, in particular for learning to drive accurately, comfortably and to drive the same way as human drivers would do.

[0013] Therefore, according to the embodiments of the present disclosure, a (desirably computer-implemented) method for training a human-like generative driving model for a vehicle is provided. The method comprises the steps of:

a – obtaining a set of video data of driving scenes performed by a human driven vehicle,

b - obtaining a data set of human driving maneuvers carried out during the driving scenes,

c – training a generative driving model using the set of video data and the data set of human driving maneuvers,

wherein the training step of c is augmented with an adversarial training scheme such that the prediction of the trained generative driving model becomes more human like.

[0014] By providing such a method, it becomes possible to take advantage of the advance of adversary learning to learn human-like driving. Specifically, a discriminator may be trained, together with the driving model, to distinguish human driving and machine driving. The driving model is trained to be accurate, comfortable, and at the same time to fool the discriminator so that it believes that the driving performed by the method was by a human driver. A new evaluation criterion is proposed to score the human-likeness of a driving model.

[0015] As a further advantage, the learning procedure is desirably improved from a pointwise prediction to a sequence-based prediction.

[0016] The generative driving model desirably outputs predicted driving maneuvers. Accordingly, the model is desirably configured to steer autonomously a vehicle based on said outputted predicted driving maneuvers. For example, the driving maneuvers may comprise any kind of maneuvers for driving control of the vehicle, e.g. simple maneuvers as steering or braking, or more complex maneuver line taking a turn by a combination of braking, steering and re-accelerating.

[0017] The adversarial training scheme may comprise:

c1 - training a discriminator model based on the predicted driving maneuvers outputted by the generative driving model and corresponding human driving maneuvers of the data set to discriminate between human and machine maneuvers, and

- 5 c2 - forcing the generative driving model to learn more human like driving (e.g. by penalizing the model using an adversary loss).

[0018] For example, the standard L1 and/or L2 loss may be augmented by an adversary loss which is based on a discriminator model trained to distinguish human driving and machine driving.

- 10 [0019] The step of obtaining the set of video data may further comprise:

a1 – obtaining a route planning data set representing route information according to which the human driven vehicle have performed the driving scenes of the set of video data,

- a2 - enriching the set of video data by the route planning data (set), such that
15 the accuracy of the predicted driving maneuvers outputted by the generative driving model is increased.

For example, the set of video data may be enriched with numerical map data from HERE Technologies.

[0020] The step of training the generative driving model may comprise:

- 20 training the generative driving model based on a predefined loss function (e.g. the L1 and/or the L2 loss) augmented by an adversary loss which is based on the output of the discriminator model.

- [0021] The generative driving model may receive as an input video data and data of human driving maneuvers of past time steps, and output predicted
25 driving maneuvers for future time steps.

[0022] In case the set of video data are enriched by the route planning data before being inputted to the generative driving model, the generative driving model may receive as a further input the vehicle location in past time steps.

- [0023] For example, given the video I, the map information M, and the
30 vehicle's location L, a deep neural network may be trained to predict the steering angle s and speed v for a future time step. All data inputs may be synchronized and sampled at the same sampling rate f , meaning the vehicle makes driving decision every $1/f$ seconds. The inputs and outputs may be represented in this discretized form.

- 35 [0024] The generative driving model may be a deep neural network.

[0025] The discriminator model may be a deep neural network.

[0026] For example, the driving model developed in S. Hecker, D. Dai, and L. Van Gool. End-to-end learning of driving models with surround-view cameras and route planners, ECCV, 2018, may be adopted.

5 [0027] The present disclosure further relates to a system for training a human-like generative driving model for a vehicle. the system comprises:
a module A for obtaining a set of video data of driving scenes performed by a human driven vehicle,
a module B for obtaining a data set of human driving maneuvers carried out
10 during the driving scenes,
a module C for training a generative driving model using the set of video data, and the data set of human driving maneuvers,
wherein training in module C is augmented with an adversarial training scheme such that the prediction of the trained generative driving model becomes more
15 human like.

[0028] The system may comprise further (sub-) modules and features corresponding to the features of the method described above.

[0029] Moreover the present disclosure relates to a system for predicting human-like driving maneuvers of a vehicle, comprising the (trained) model of
20 step c or of module C, as descried above.

[0030] Furthermore the present disclosure relates to a computer program including instructions for executing the steps of a method, as described above, when said program is executed by a computer.

[0031] This program can use any programming language and take the form
25 of source code, object code or a code intermediate between source code and object code, such as a partially compiled form, or any other desirable form.

[0032] Finally, the present disclosure relates to a recording medium readable by a computer and having recorded thereon a computer program including instructions for executing the steps of a method, as described above.

30 [0033] The information medium can be any entity or device capable of storing the program. For example, the medium can include storage means such as a ROM, for example a CD ROM or a microelectronic circuit ROM, or magnetic storage means, for example a diskette (floppy disk) or a hard disk.

[0034] Alternatively, the information medium can be an integrated circuit in which the program is incorporated, the circuit being adapted to execute the method in question or to be used in its execution.

5 [0035] It is intended that combinations of the above-described elements and those within the specification may be made, except where otherwise contradictory.

[0036] It is to be understood that both the foregoing general description and the following detailed description are exemplary and explanatory only and are not restrictive of the disclosure, as claimed.

10 [0037] The accompanying drawings, which are incorporated in and constitute a part of this specification, illustrate embodiments of the disclosure and together with the description, and serve to explain the principles thereof.

BRIEF DESCRIPTION OF THE DRAWINGS

15 [0038] Fig. 1 shows a schematic flow chart of the steps of a method for training a human-like generative driving model according to embodiments of the present disclosure; and

[0039] Fig. 2 shows a schematic block diagram of a system according to embodiments of the present disclosure.

20

DESCRIPTION OF THE EMBODIMENTS

[0040] Reference will now be made in detail to exemplary embodiments of the disclosure, examples of which are illustrated in the accompanying drawings. Wherever possible, the same reference numbers will be used throughout the
25 drawings to refer to the same or like parts.

[0041] End-to-end driving allows developing promising driving models based on camera data, cf. e.g.:
the driving model developed in S. Hecker, D. Dai, and L. Van Gool. End-to-end learning of driving models with surround-view cameras and route planners,
30 ECCV, 2018.

[0042] The focus has mainly been though on perception, not so much navigation. Thus far, the representations for navigation are either primitive directional commands in a simulation environment or rendered videos of planned routes in real-world environments.

[0043] Fig. 1 shows a schematic flow chart of the steps of a method for training a human-like generative driving model according to embodiments of the present disclosure. For example, the driving model developed in the publication cited above (S. Hecker et. al., 2018) may be adopted in the present disclosure.

[0044] In particular, the used core model may consist of a fine-tuned Resnet34 CNN (cf. e.g. K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016) to process sequences of front facing camera images, followed by two regression networks to predict steering wheel angle and vehicle speed. The architecture may thus be similar to the baseline model from the publication cited above (S. Hecker et. al., 2018).

[0045] In a first step S01 a set of video data of driving scenes performed by a human driven vehicle is obtained. The set of video data may be e.g. the Drive360 dataset, as described in the publication cited above (S. Hecker et. al., 2018).

[0046] In an optional step S01a (not shown) a route planning data set representing route information according to which the human driven vehicle have performed the driving scenes of the set of video data may be additionally obtained, e.g. map data from Here Technologies.

[0047] In a further optional step S01b (not shown) the set of video data may be enriched by the route planning data, such that the accuracy of the predicted driving maneuvers outputted by the generative driving model is increased.

[0048] It is hence proposed by the present disclosure to (1) augment real-world driving data with numerical map data from e.g. HERE Technologies; and (2) design map features believed to be relevant for driving and integrating them into a driving model.

[0049] In a further step S02 (which may be carried out before, after or at the same time as step S01) a data set of human driving maneuvers carried out during the driving scenes is obtained.

[0050] Accordingly, said data set of human driving maneuvers is a further input for the model providing information regarding the driving control gestures of humans when driving the vehicles in the driving scenes of the set of video data.

[0051] In a subsequent step S03 a generative driving model is trained using the set of video data and the data set of human driving maneuvers. Said training is augmented with an adversarial training scheme such that the prediction of the trained generative driving model becomes more human like.

- 5 [0052] In particular, in an optional step S03a (not shown) a discriminator model may be trained based on the predicted driving maneuvers outputted by the generative driving model and corresponding human driving maneuvers of the data set to discriminate between human and machine maneuvers. The generative driving model may be forced to learn more human like driving in a
10 further optional step S03b (not shown).

- [0053] For example, given the video I , the map information M , and the vehicle's location L , a deep neural network is trained to predict the steering angle s and speed v for a future time step. All data inputs are synchronized and sampled at the same sampling rate f , meaning the vehicle makes driving
15 decision every $1/f$ seconds. The inputs and outputs are represented in this discretized form. It is used t to indicate the time stamp, such that all data can be indexed over time. For example, I_t indicates the current video frame and v_t the vehicle's current speed. Similarly, I_{t-k} is the k^{th} previous video frame and s_{t-k} is the k^{th} previous steering angle. Since predictions need to rely on data of
20 previous time steps, the k recent video frames are denoted by $I_{[t-k+1,t]} \equiv \langle I_{t-k+1}, \dots, I_t \rangle$, and the k recent map representations by $M_{[t-k+1,t]} \equiv \langle M_{t-k+1}, \dots, M_t \rangle$. The goal is to train a deep network that predicts desired driving actions from the visual observations and the planned route. The learning task can be defined as:

25

$$F: (\mathcal{I}_{[t-k+1,t]}, \mathcal{M}_{[t-k+1,t]}) \rightarrow \mathcal{S}_{t+1} \times \mathcal{V}_{t+1} \quad (1)$$

- where $\mathcal{S}_{t+1}$ represents the steering angle space and $\mathcal{V}_{t+1}$ the speed space for future time $t+1$. $\mathcal{S}$ and $\mathcal{V}$ can be defined at several levels of granularity. The continuous values directly recorded from the car's CAN bus may be considered,
30 where $\mathcal{V} = \{V | 0 \leq V \leq 180 \text{ for speed and } \mathcal{S} = \{S | -720 \leq S \leq 720\} \text{ for steering angle in this case. Here, kilometer per hour (km/h) is the unit of } v, \text{ and degree } (^{\circ}) \text{ the unit of } s. M_t \text{ is either a rendered video frame from the TomTom route planner (cf. S. Hecker et. al., 2018), or the engineered features for the}$

numerical maps from HERE Technologies (as described below), or the combination of both.

[0054] In order to keep notations concise, the synchronized data (I, M) may be denoted as D. Without loss of generality, the training data are assumed to consist of a long sequence of driving data with T frames in total. Then the basic driving model is to learn the prediction function for the steering angle

$$\hat{s}_{t+1} \leftarrow f_{st}(\mathcal{D}_{[t-k+1,t]}), \quad (2)$$

and the velocity

$$\hat{v}_{t+1} \leftarrow f_{ve}(\mathcal{D}_{[t-k+1,t]}), \quad (3)$$

with the objective

$$\sum_{t=1}^{T-1} (|\hat{s}_{t+1} - s_{t+1}| + \lambda |\hat{v}_{t+1} - v_{t+1}|), \quad (4)$$

where $\hat{s}$ and $\hat{v}$ are predicted values, and s and v are the ground truth values.

[0055] The learning under Eq. 4 is straightforward and can be implemented with any standard deep network. This objective, however, assumes the driving decisions at each time step are independent from each other. It is believed that this may be an over-simplification because driving decisions indeed exhibit strong temporal dependencies within a relatively short time range. In the following section, the objective according to the present disclosure is reformulated by introducing ride comfort and human-likeness score to better model the temporal dependency of driving actions.

Accurate and Comfortable Driving

[0056] Multiple concepts relating to driving comfort have been proposed and discussed, such as apparent safety, motion sickness, level of controllability and resulting force. While those are all relevant, some are hard to quantify. It is hence chosen to reduce motion sickness, which has been shown to be largely caused by the vehicle's longitudinal and lateral oscillations, cf. e.g.:

G. M. Turner M1. Motion sickness in public road transport: the effect of driver, route and vehicle. *Ergonomics*, (1646-64), 1999.

[0057] Due to the short-term predictive nature of most end-to-end driving models, substantial jerking is an inherent problem. The used comfort component aims at reducing jerk by imposing a temporal smoothness constraint on the longitudinal and lateral oscillations, by minimizing the second derivative of consecutive steering angle and speed predictions.

[0058] Before introducing ride comfort and human-like driving, Eq. 4 is reformulated. If the number of consecutive predictions that need to be optimized jointly is denoted by O , then minimizing Eq. 4 is equivalent to minimizing

$$\sum_{t=1}^{T-O} \sum_{o=1}^O (|\hat{s}_{t+o} - s_{t+o}| + \lambda |\hat{v}_{t+o} - v_{t+o}|). \quad (5)$$

[0059] Then for every O consecutive frames starting at time t , the loss of driving accuracy will be

$$\mathcal{L}_t^{\text{acc}} = \sum_{o=1}^O (|\hat{s}_{t+o} - s_{t+o}| + \lambda |\hat{v}_{t+o} - v_{t+o}|) \quad (6)$$

[0060] Now the objective function can be presented for accurate and comfortable driving as

$$\sum_{t=1}^{T-O} \sum_{o=1}^O (\mathcal{L}_t^{\text{acc}} + \zeta_1 \mathcal{L}_t^{\text{com}}), \quad (7)$$

where

$$\begin{aligned} \mathcal{L}_t^{\text{com}} = & \sum_{o=2}^{O-1} \frac{|\hat{s}_{t+o-1} - 2\hat{s}_{t+o} + \hat{s}_{t+o+1}|}{(1/f)^2} \\ & + \lambda \frac{|\hat{v}_{t+o-1} - 2\hat{v}_{t+o} + \hat{v}_{t+o+1}|}{(1/f)^2}. \end{aligned} \quad (8)$$

[0061] ζ_1 is a trade-off parameter to balance the two costs. By optimizing under the objective in Eq. 7, consecutive predictions are learned and optimized together for accurate and comfortable driving.

Accurate, Comfortable & Human-like Driving

[0061] If autonomous cars behave differently from human-driven cars, it is hard for humans to predict their future actions. This unpredictability can cause accidents. Thus, it is argued that it is important to design human-like driving algorithms from the very start. Hence, a human-likeness score is introduced. The higher the value, the closer to human driving. Since it is hard to manually define what a human driving style is – as was done for general comfort measures, adversarial learning is adopted to model it.

[0062] An adversarial learning method consists of a generator and discriminator. The driving model of the present disclosure as defined in Eq. 4 or in Eq. 7 is the generator G . Now the training objective for the discriminator will be described. For convenience, the short trajectories of O frames described above are named as drivelets. Given the outputs of the driving model for a drivelet $\hat{B}_t = (\hat{s}_{t+1}, \dots, \hat{s}_{t+O}, \hat{v}_{t+1}, \dots, \hat{v}_{t+O})$ and its corresponding ground truth from the human driver $B_t = (s_{t+1}, \dots, s_{t+O}, v_{t+1}, \dots, v_{t+O})$, the goal is to train a fully-connected discriminator D using the cross-entropy loss to classify the two classes (i.e. machine and human).

[0063] The drivelet at t is forwarded to G to obtain the driving actions $\hat{B}_t$. To make autonomous driving more human-like is equivalent to letting the distribution of $\hat{B}_t$ approximate that of B_t . Thus the loss for human-like driving according to the present disclosure is defined as an adversarial loss:

$$\mathcal{L}_t^{\text{hum}} = -\log(D(\hat{B}_t)^1), \quad (9)$$

where $D(\hat{B}_t)^1$ is the probability of classifying $\hat{B}_t$ as human driving.

[0064] Putting everything together, the objective for accurate, comfortable and human-like driving according to the present disclosure is as follows:

$$\mathcal{L} = \sum_{t=1}^{T-O} \sum_{o=1}^O (\mathcal{L}_t^{\text{acc}} + \zeta_1 \mathcal{L}_t^{\text{com}}) + \zeta_2 \mathcal{L}_t^{\text{hum}}. \quad (10)$$

ζ_2 is a trade-off parameter to control the contributions of the costs. In keeping with adversarial learning, the training is conducted under the following min-max criterion:

$$\max_G \min_D \mathcal{L}(I, M). \quad (11)$$

Obtaining HERE Map Data

5

[0065] The set of video data according to the present disclosure may be provided by panoramic videos recorded by vehicle comprising one or several cameras oriented toward the environment in at least one of the front, the sides, and the back of the vehicle, e.g. by the Drive360 video data set. Drive360 features 60 hours of real-world driving data over 3000 km. The Drive360 is e.g. augmented with HERE Technologies map data. Drive360 offers a time stamped GPS trace for each route recorded. A path-matcher is used based on a hidden markov model employing the Viterbi algorithm (cf. e.g. G. D. Forney. The viterbi algorithm. Proceedings of the IEEE, 61(3):268–278, 1973) to calculate the most likely path traveled by the vehicle during dataset recording, snapping the GPS trace to the underlying road network. This improves the localization accuracy significantly, especially in urban environments where the GPS signal may be weak and noisy. Through the path matcher a map matched GPS coordinate is obtained for each time stamp, which is then used to query the HERE Technologies map database to obtain the various types of navigation data.

[0066] Following S. Hecker et. al., 2018, a fine-tuned AlexNet (cf. e.g. A. Krizhevsky, I. Sutskever, and G. E. Hinton. Imagenet classification with deep convolutional neural networks. In Advances in Neural Information Processing Systems. 2012) may be used to process the visual map representation from the TomTom Go App.

[0067] For learning comfortable driving, no extra network is needed. The respective loss may computed according to Eq. 7 and gradients are back propagated to adjust the driving network. In order to learn human-like driving, a fully-connected, three-layer discriminator network may be used to model human-like driving. The loss may be computed according to Eq. 9 to adjust the driving network.

[0068] Fig. 2 shows a schematic block diagram of a system according to embodiments of the present disclosure.

[0069] In this figure, a system 200 for training a model has been represented. This system 200, which may be a computer, comprises a processor 201 and a non-volatile memory 202. The system 200 may also comprise, be configured to be integrated in or form a part of a vehicle 400. The system 200 may not only be configured for training a human-like generative driving model for a vehicle but also to apply the trained model to autonomously drive a vehicle (in particular in case it is part of a vehicle 400).

[0070] The system 200 may further be connected to a (passive) optical sensor 300, in particular a digital camera (e.g. integrated into the vehicle and being oriented to at least one of the front, the sides and the back). The digital camera 300 is configured such that it can record a scene in front of the vehicle 400, and in particular output digital data providing appearance (color, e.g. RGB) information of the scene. The camera 300 desirably generates image data comprising a 2D or 3D image of the environment. There may also be provided a set of monocular cameras which generate a panoramic 2D or 3D image. The output of the camera 300 may be used as video data of driving scenes for training the model (cf. step S01 of the method described above) and/or as input for a trained model, based on which the trained model autonomously controls driving the vehicle.

[0071] In the non-volatile memory 202, a set of instructions is stored and this set of instructions comprises instructions to perform a method for training a model.

[0072] In particular, these instructions and the processor 201 may respectively form a plurality of modules:

- a module A for obtaining (S01) a set of video data of driving scenes performed by a human driven vehicle,
 - a module B for obtaining a data set of human driving maneuvers carried out during the driving scenes,
 - a module C for training (S03) a generative driving model using the set of video data, and the data set of human driving maneuvers,
- wherein training in module C is augmented with an adversarial training scheme such that the prediction of the trained generative driving model becomes more human like.

[0073] Throughout the description, including the claims, the term "comprising a" should be understood as being synonymous with "comprising at

least one" unless otherwise stated. In addition, any range set forth in the description, including the claims should be understood as including its end value(s) unless otherwise stated. Specific values for described elements should be understood to be within accepted manufacturing or industry tolerances
5 known to one of skill in the art, and any use of the terms "substantially" and/or "approximately" and/or "generally" should be understood to mean falling within such accepted tolerances.

[0074] Although the present disclosure herein has been described with reference to particular embodiments, it is to be understood that these
10 embodiments are merely illustrative of the principles and applications of the present disclosure.

[0075] It is intended that the specification and examples be considered as exemplary only, with a true scope of the disclosure being indicated by the following claims.

CLAIMS

1. A method for training a human-like generative driving model for a vehicle, comprising the steps of:
 - 5 a – obtaining (S01) a set of video data of driving scenes performed by a human driven vehicle,
 - b - obtaining (S02) a data set of human driving maneuvers carried out during the driving scenes,
 - c – training (S03) a generative driving model using the set of video data and
10 the data set of human driving maneuvers,
wherein the training step (S03) of c is augmented with an adversarial training scheme such that the prediction of the trained generative driving model becomes more human like.
- 15 2. The method according to the claim 1, wherein the generative driving model outputs predicted driving maneuvers.
3. The method according to claim 1 or 2, wherein the adversarial training scheme comprises:
 - 20 c1 – training (S03a) a discriminator model based on the predicted driving maneuvers outputted by the generative driving model and corresponding human driving maneuvers of the data set to discriminate between human and machine maneuvers, and
 - c2 - forcing (S03b) the generative driving model to learn more human like
25 driving.
4. The method according to any one of the preceding claims, wherein the step of obtaining (S01) the set of video data further comprises:
 - a1 – obtaining (S01a) a route planning data set representing route information
30 according to which the human driven vehicle have performed the driving scenes of the set of video data,
 - a2 - enriching (S01b) the set of video data by the route planning data, such that the accuracy of the predicted driving maneuvers outputted by the generative driving model is increased.

5. The method according to any one of the preceding claims, wherein the step of training (S03) the generative driving model comprises: training the generative driving model based on a predefined loss function augmented by an adversary loss which is based on the output of the discriminator model.
6. The method according to any one of the preceding claims, wherein the generative driving model receives as an input video data and data of human driving maneuvers of past time steps, and outputs predicted driving maneuvers for future time steps.
7. The method according to the preceding claim, wherein in case the set of video data are enriched by the route planning data before being inputted to the generative driving model, the generative driving model receives as a further input the vehicle location in past time steps.
8. The method according to the preceding claim, wherein the generative driving model is a deep neural network, and/or the discriminator model is a deep neural network.
9. A system for training a human-like generative driving model for a vehicle, comprising:
a module A for obtaining a set of video data of driving scenes performed by a human driven vehicle,
a module B for obtaining a data set of human driving maneuvers carried out during the driving scenes,
a module C for training a generative driving model using the set of video data, and the data set of human driving maneuvers,
wherein training in module C is augmented with an adversarial training scheme such that the prediction of the trained generative driving model becomes more human like.
10. A system for predicting human-like driving maneuvers of a vehicle, comprising the model of step c of any one of claims 1 to 8 or of module C of claim 9.

11. A computer program including instructions for executing the steps of a method according to any one of claims 1 to 8 when said program is executed by a computer.

5

12. A recording medium readable by a computer and having recorded thereon a computer program including instructions for executing the steps of a method according to any one of claims 1 to 8.

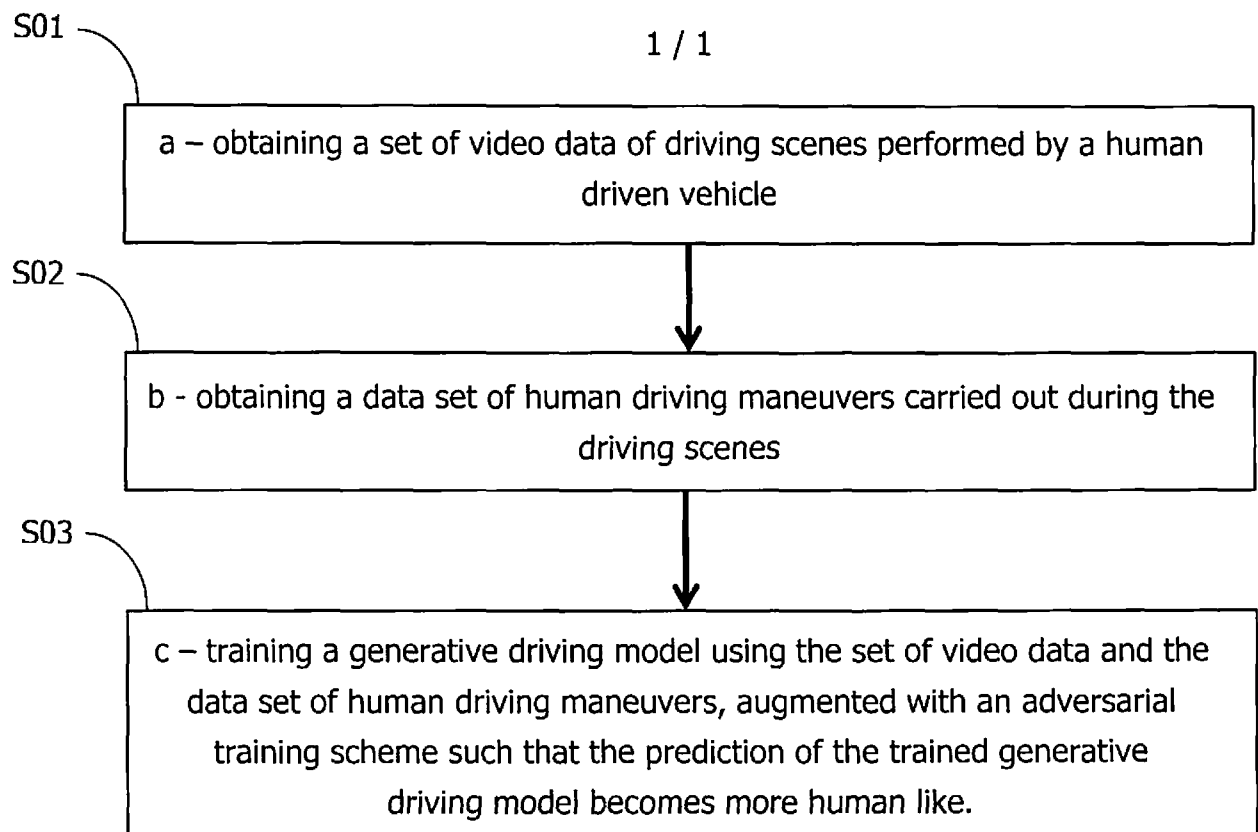


Fig. 1

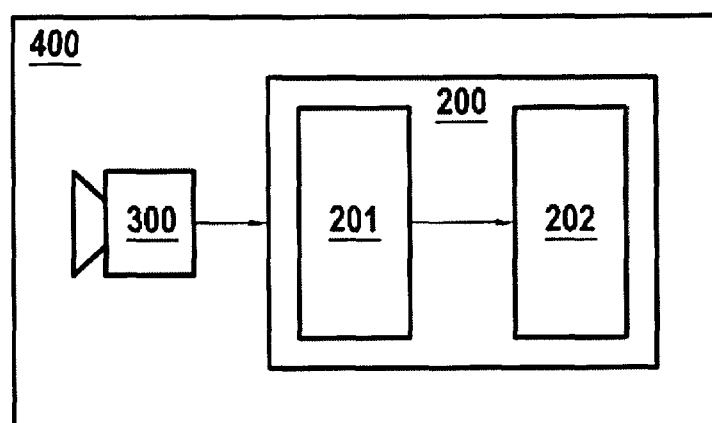


Fig. 2

INTERNATIONAL SEARCH REPORT

International application No
PCT/EP2019/055786

A. CLASSIFICATION OF SUBJECT MATTER
INV. G06K9/00 G06K9/62
ADD.

According to International Patent Classification (IPC) or to both national classification and IPC

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)
G06K

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

EPO-Internal, WPI Data

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
X	LUONA YANG ET AL: "Real-to-Virtual Domain Unification for End-to-End Autonomous Driving", ARXIV.ORG, CORNELL UNIVERSITY LIBRARY, 201 OLIN LIBRARY CORNELL UNIVERSITY ITHACA, NY 14853, 10 January 2018 (2018-01-10), XP081195420, figures 2 and 3, sections 1, 3.1 and 4.1; equations 1-4	1,2,5, 9-12
Y	----- -/--	4,6-8



Further documents are listed in the continuation of Box C.



See patent family annex.

* Special categories of cited documents :

"A" document defining the general state of the art which is not considered to be of particular relevance

"E" earlier application or patent but published on or after the international filing date

"L" document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified)

"O" document referring to an oral disclosure, use, exhibition or other means

"P" document published prior to the international filing date but later than the priority date claimed

"T" later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention

"X" document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone

"Y" document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art

"&" document member of the same patent family

Date of the actual completion of the international search

10 January 2020

Date of mailing of the international search report

23/01/2020

Name and mailing address of the ISA/

European Patent Office, P.B. 5818 Patentlaan 2
NL - 2280 HV Rijswijk
Tel. (+31-70) 340-2040,
Fax: (+31-70) 340-3016

Authorized officer

Loza, Artur

INTERNATIONAL SEARCH REPORT

International application No
PCT/EP2019/055786

C(Continuation). DOCUMENTS CONSIDERED TO BE RELEVANT		
Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	Jonathan Ho ET AL: "Generative Adversarial Imitation Learning", 10 June 2016 (2016-06-10), XP055639480, Retrieved from the Internet: URL:https://papers.nips.cc/paper/6391-generative-adversarial-imitation-learning.pdf [retrieved on 2019-11-12] the whole document	1-3,5, 9-12
Y	US 2018/348763 A1 (JIANG YIFEI [US] ET AL) 6 December 2018 (2018-12-06) paragraph 30 and 37-38 and figures 1 and 4	4,8
Y	US 2007/136040 A1 (TATE EDWARD D JR [US]) 14 June 2007 (2007-06-14) paragraphs 4 and 34	6,7

INTERNATIONAL SEARCH REPORT

International application No.
PCT/EP2019/055786

Box No. II Observations where certain claims were found unsearchable (Continuation of item 2 of first sheet)

This international search report has not been established in respect of certain claims under Article 17(2)(a) for the following reasons:

1. ☐ Claims Nos.:
because they relate to subject matter not required to be searched by this Authority, namely:
2. ☐ Claims Nos.:
because they relate to parts of the international application that do not comply with the prescribed requirements to such an extent that no meaningful international search can be carried out, specifically:
3. ☐ Claims Nos.:
because they are dependent claims and are not drafted in accordance with the second and third sentences of Rule 6.4(a).

Box No. III Observations where unity of invention is lacking (Continuation of item 3 of first sheet)

This International Searching Authority found multiple inventions in this international application, as follows:

see additional sheet

1. ☒ As all required additional search fees were timely paid by the applicant, this international search report covers all searchable claims.
2. ☐ As all searchable claims could be searched without effort justifying an additional fees, this Authority did not invite payment of additional fees.
3. ☐ As only some of the required additional search fees were timely paid by the applicant, this international search report covers only those claims for which fees were paid, specifically claims Nos.:
4. ☐ No required additional search fees were timely paid by the applicant. Consequently, this international search report is restricted to the invention first mentioned in the claims; it is covered by claims Nos.:

Remark on Protest

- ☐ The additional search fees were accompanied by the applicant's protest and, where applicable, the payment of a protest fee.
- ☐ The additional search fees were accompanied by the applicant's protest but the applicable protest fee was not paid within the time limit specified in the invitation.
- ☒ No protest accompanied the payment of additional search fees.

FURTHER INFORMATION CONTINUED FROM PCT/ISA/ 210

This International Searching Authority found multiple (groups of) inventions in this international application, as follows:

1. claims: 1-3, 5, 9(completely); 10-12(partially)
discriminating between human and machine maneuvers

2. claims: 4, 6-8(completely); 10-12(partially)
enriching video data with route planning data

INTERNATIONAL SEARCH REPORT

Information on patent family members

International application No

PCT/EP2019/055786

Patent document cited in search report	Publication date	Patent family member(s)	Publication date
US 2018348763 A1	06-12-2018	CN 108974009 A US 2018348763 A1	11-12-2018 06-12-2018
US 2007136040 A1	14-06-2007	CN 1983240 A DE 102006058588 A1 US 2007136040 A1	20-06-2007 19-07-2007 14-06-2007