

Complex self-driving behaviors emerging from affordance competition in layered control architectures[☆]

Mauro Da Lio, Antonello Cherubini^{*}, Gastone Pietro Rosati Papini, Alice Plebe

Department of Industrial Engineering, University of Trento, Via Sommarive 9, Povo, TN, Italy

ARTICLE INFO

Keywords:

Self-driving car
Autonomous driving
Lane merge
Motorway intersection
Emergent behavior
Cognitive inspiration

ABSTRACT

The deployment of autonomous driving technology is hindered by “corner cases”: unusual nuanced conditions that the self-driving software cannot understand and act fully. We argue that some corner cases originate from a “narrow AI” approach, which lacks the general knowledge that humans exploit when dealing with these cases. We propose an alternative that can be seen as a step toward features of Artificial General Intelligence.

We exploit the biological principle of affordance competition in layered control architectures to create an artificial agent that realizes emergent, adaptive, and logical behaviors without programming case-specific rules or algorithms. We give six different examples of simple and complex emergent behaviors. For the case study of merge scenarios, we contrast the approach of this paper with an algorithmic solution of the literature. The ideas presented here (if not the whole agent’s sensorimotor organization) could be used to improve the robustness and flexibility of self-driving technology.

1. Introduction

The notion of self-driving vehicles is old-fashioned, but is still an active field of industrial research. A CB Insight report (CB Insights, 2015–2020) listed more than 40 companies that work on autonomous vehicles in 2020.

The main motivations for automating driving are:

- (1) Safety. Under the assumption that humans are bad drivers and that it is possible to automate driving.
- (2) Providing new mobility services to reduce traffic congestion, energy consumption and pollution, and for people who cannot drive.
- (3) Maintaining technological and market leadership.

For decades, the achievement of self-driving vehicles has been a mirage; it has become closer recently (but not without issues), in parallel with the rapid advancement of Artificial Intelligence (AI) powered by deep neural networks.

1.1. Autonomous vehicles within AI

It is natural to ask what role AI plays in the research on autonomous vehicles (AV); more specifically, can AV technology be considered part

of Artificial General Intelligence (AGI)? At first glance, the answer appears to be negative. According to Ben Goertzel, who introduced the concept of AGI (Goertzel, 2014; Goertzel, Iklé, & Wigmore, 2012; Goertzel & Pennachin, 2006), the main characteristic of AGI is the ability to perform a variety of tasks in many different contexts and environments. AGI differs from the so-called “narrow AI” research, which aims to create intelligent artifacts specialized in one task. From this point of view, autonomous driving appears to fall precisely within the field of narrow AI as it pursues a single specialized task: driving a vehicle.

In this paper, we argue that the main limitations of mainstream research on AV come from its account as narrow AI. The idea of developing AV as an application of narrow AI makes it almost impossible to attain a higher level of driving autonomy.

Our claim is supported by the weakness of the assumption at point 1 of the above list. In contrast to shallow perception, human beings are excellent drivers. The fatality rate of an “average” driver in the United States or the EU is one fatality per 100 million miles driven and one severe accident every 12 million miles (Kalra & Paddock, 2016). The 40.000 deaths per year in the EU and 33.000 in the US, if not seen with the vast driven distances in the nations, may cause the false perception that humans are bad drivers. If automated vehicles were to be designed

[☆] This work was supported by the European Commission with Grant H2020 731593 (Dreams4Cars) and by the EU funded program FSE REACT EU via the Italian Ministry of Education (D.M.1062).

^{*} Corresponding author.

E-mail addresses: mauro.dalio@unitn.it (M. Da Lio), antonello.cherubini@unitn.it (A. Cherubini), gastone.rosatipapini@unitn.it (G.P. Rosati Papini), alice.plebe@unitn.it (A. Plebe).

<https://doi.org/10.1016/j.cogsys.2022.12.007>

Received 13 October 2022; Received in revised form 3 December 2022; Accepted 12 December 2022

to meet the safety goal (point 1), it would be clear that these systems should be significantly more reliable than human drivers. That would raise the bar to the level of one death every 1–10 billion miles (which still means about 400 deaths per year in the EU) or even more.

However, the recent development of autonomous driving technology appears to be impeded by the continued emergence of many edge and corner cases (Adams, 2020; Anderson, 2020; Arieff, 2021; Eliot, 2021; The New York Times, 2021a, 2021b). The prototypes of driverless vehicles show a long tail of unexpected situations in which cars are unable to act (Boggs, Arvin, & Khattak, 2020; Dixit, Chand, & Nair, 2016; Lv et al., 2018), or act dangerously (Marshall & Davies, 2019; US National Transport Safety Board, 2019).

The difference between narrow and general intelligence becomes crucial when dealing with edge- and corner-cases. During unexpected situations and emergencies, the most appropriate reaction might come from a broad knowledge of objects in the world and intuitive cognition of the physics that rules their motion and interactions (we will give an example in Section 2.1). However, note that there is no interest in developing AV capable of performing other tasks in addition to driving. For this reason, it may never fully belong to a strict definition of AGI. Still, we believe that advanced autonomous driving systems should be at an intermediate point in the intelligence spectrum between narrowly specialized AI models and AGI.

1.2. Paper contribution

The work presented here attempts to move from narrow AI to AGI by using certain architectural principles belonging to AGI, with the objective of better coping with driving situations requiring high-level intelligence.

The industrial approach to dealing with corner cases relies on large-scale research and development programs (Connected Automated Driving, 2021) that combine experiments, simulation, and standardization activities; for example, the PEGASUS (Project PEGASUS, 2019) and HEADSTART (EU project HEADSTART, 2019) projects in the EU.

This paper presents an agent’s sensorimotor architecture that produces polite behavior as an emergent feature of the sensorimotor system. We speculate that such a property may ease the burden of dealing with corner cases.

Among the various approaches to AGI (Goertzel, 2014), the agent presented here loosely embraces three of them. First, the agent takes inspiration from several features of brain organization regarding perception and action selection. Second, we are in agreement with the “emergentist approach”, in which internal representations of concepts and behavioral schemes emerge from lower-level dynamics. Third, our agent follows the embodied perspective in which intelligence is something that physical moving agents do in physical environments.

The novel contribution is found in Section 4. It is preceded by a summary of the agent’s architecture in Section 3, which was published previously, but is summarized here to the extent necessary to explain the contribution of Section 4.

1.3. Paper organization

In the next Section, we review the two main approaches to AV that follow the narrow AI account: in Section 2.1, the approach based on an engineering decomposition of the overall system into sequential modules (e.g., sense–think–act); in Section 2.2, the opposite approach known as *end-to-end*. We argue that both are brittle and harbor seeds for corner cases.

In Section 2.3 we present the layered control architecture, a biological solution capable of generating emergent adaptive behavior with successful robotics applications. We clarify how our layered control architecture with affordance competition works in Section 3. Finally, Section 4 constitutes the core of the paper. It gives six different examples of emergent behaviors, in most cases for non-trivial scenarios (when possible, recalling examples of the same scenarios handled in a traditional way). Multimedia materials support the demonstrations.

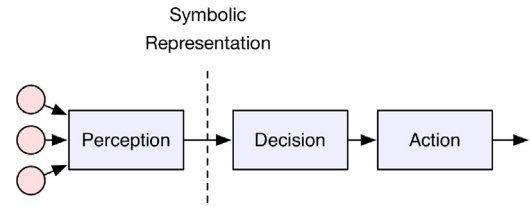


Fig. 1. The sense–think–act architecture.

2. State-of-the-art of self-driving architectures

2.1. The sense–think–act architecture

The architecture almost universally used for self-driving software is the sense–think–act shown in Fig. 1 (Claussmann, Revilloud, Gruyer, & Glaser, 2019; González, Pérez, Milanés, & Nashashibi, 2016; Karakazas, Quddus, Chen, & Deka, 2015; Paden, Čáp, Yong, Yershov, & Frazzoli, 2016; Pendleton et al., 2017). Interestingly, as a model for human intelligence, it is also known as the Cartesian model, but its ability to explain natural cognition is questioned, for example, Barsalou (2008), Brooks (1991), Keijzer (2002), Seitz (2000), Van Gelder (1995).

From an engineering perspective, the scheme of Fig. 1 is a convenient decomposition of a system’s functions, which works fine when the system states can be modeled perfectly beforehand (which is the case of many engineered systems).

However, it shows weaknesses when faced with situations that cannot be perfectly specified in advance. In this model, a designer decides how to “represent” the world and, in particular, the symbols used for that representation (for example, the classes of objects, such as cars, pedestrians, etc.). These symbols are abstract labels (Keijzer, 2002) that do not contain any information about how real things work. The self-driving behaviors — what to do with given classes — are programmed by a human designer.

The reasons why we argue that the scheme of Fig. 1, as an artificial cognitive system, is brittle are at least twofold: (1) the inherent deficiency of narrow-AI and (2) the complexity of programming every behavior.

2.1.1. Inherent narrow-AI deficiency

Fig. 2 gives an example of a typical outcome of a narrow-AI perceptual system that turns out to lack important information. It shows the output of a state-of-the-art semantic segmentation network, where the object labels lack any inferential meaning representation, unlike an AGI system.

The front minivan is classified as a vehicle not different from the others. Information about the dangerousness of the minivan load, which could fall onto the road through the open door, is lost. The reason is that the neural network’s output symbols have been predefined without considering this unusual situation (and many others). Even if we could retrain the network with one additional class for the special vehicle (which would be a considerable undertaking), a quick Internet search reveals an infinite number of hanging-load situations, some requiring urgent action and others with little danger. Therefore, the class in itself still says nothing about which behavior is required: the meaning of the symbols and, consequently, action planning, is left to a human programmer with limited cues about what to do for virtually infinite variants of unforeseen situations.

2.1.2. Complexity of behavior programming

Even if the output coding of the perceptual system might include some meaningful content of objects, decision-making can still be problematic. Behavior selection and optimization are typically non-convex, and many special maneuvers are often handled by systems of rules or *ad hoc* algorithms. For example, negotiation in a merge scenario is solved



Fig. 2. A driving scenario in which a vehicle with a dangerously open door is confused with the others.

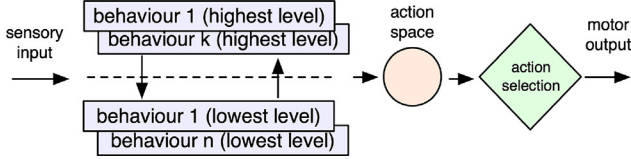


Fig. 3. Layered control architecture.

with a specific algorithm (which requires the creation of a phantom obstacle) in Kreutz and Eggert (2021). For the same courtesy merge example, the complexity related to behavior generation may also be glimpsed in a recent Ph.D. thesis (Menéndez-Romero, 2021), which shows how the function breaks down into a plethora of subcases.

2.2. End-to-end architectures

A recent alternative to the traditional sequential decomposition of autonomous driving tasks lies in the extreme opposite. End-to-end learning approaches train a (deep) neural network controller by giving expert driver demonstration examples. A notable demonstration by NVIDIA is the lane keeping function (Bojarski et al., 2016), where a deep neural network realizes a sensorimotor loop that produces control commands in response to raw camera data. The network learns patterns that could be considered as symbols created on purpose, not a rigid set of symbols like in the sense–think–act paradigm (Bojarski et al., 2017).

However, because of the narrow AI scope, the most appealing feature of the end-to-end strategy — dispensing with programmed internal representations and algorithms — is also a major source of trouble. It is impossible to acquire an implicit knowledge of objects and the world phenomena from a single narrow task of driving. The end-to-end approach struggles with the vastness and diversity of the training set necessary “to train a generalizable model which can drive in all different environments” (Bansal & Ogale, 2018), which is related to another recently discovered issue: the so-called causal confusion, i.e. the inability to grasp the “causal structure of the interaction between the expert and the environment” (De Haan, Jayaraman, & Levine, 2019).

Finally, human actions are a superintended choice between different affordances (Marti, Morice, & Montagne, 2015; Pezzulo & Cisek, 2016). This level of explanation is lacking with end-to-end approaches: it remains unclear which alternative actions have been evaluated and why they have been discarded.

2.3. Layered control architectures

Layered control architectures are shown in Fig. 3. They combine a subsumption architecture (Brooks, 1986) (the violet stack on the left) with the affordance competition principle (Cisek, 2007; Prescott, Redgrave, & Gurney, 1999).

At any given time, the largest number of potential actions that an agent detects in the environment — also called *affordances* (Gibson, 1986) — are simultaneously primed by the violet stack, creating a large pool of opportunities (the light orange circle).

The optimal action is then selected from the pool (action selection) through a robust centralized *competition* process (Cisek, 2007). The pool of actions instantiated before the final choice explains which alternative affordances were considered.

3. Our implementation of layered control

We have realized a self-driving agent with a layered control architecture, which is described in detail in Da Lio, Donà, Rosati Papini, and Gurney (2020).

For the reader’s convenience, we summarize its operation here. However, we remark that this summary should not be considered a complete self-contained description of the agent. This description is limited to what is necessary to understand the novel contributions of the paper, which will be presented in Section 4. The Appendix A gives further details linked to the previous literature.

Technically speaking, the layers of the subsumption stack are “inverse models” that link the sensory effect to the motor command that causes it. In principle, they can be learned via generalized motor babbling, i.e., by producing exploratory motor commands and observing the sensory effects. Examples of bootstrapping subsumption architectures in this way were studied in the European projects COSPAL (EU project COSPAL, 2007) and DIPLECS (EU project DIPLECS, 2010). In biology, learning cerebellar inverse models from sensory-motor pairs is explained in Porrill, Dean, and Anderson (2012).

In driving environments, testing actions in the real world may be dangerous. In the EU Dreams4Cars project (EU project Dreams4Cars, 2019), we used an indirect approach that synthesizes the inverse models via “mental simulation” processes (Da Lio, Donà, Rosati Papini, Biral, & Svensson, 2020a), that is, interacting in a safe sandbox created with pre-learned forward models (the agent’s offline operation). The process is inspired by human mental synthesis that occurs in various mental states, including when thinking about new actions and during sleep and dreams.

3.1. Online operation

During action (the agent’s online operation), the inverse models in the subsumption stack respond to different affordances offered by the environment. For example, in Fig. 4, the sensorimotor model labeled “lane follow” reacts to the affordance corresponding to traveling in the lane (a_1). The inverse model labeled “lane change left” reacts to the possible action of traveling in both the current and the left lane (a_2). Each inverse model estimates the salience $Q_i(r_0, j_0)$ for affordance a_i as:

$$Q_i(r_0, j_0) = \max_{r(t), j(t)} \mathbb{E} [R_i[r(t), j(t)]],$$

$$s.t.$$

$$r(0) = r_0, j(0) = j_0,$$

where $r(t), j(t)$ are the lateral and longitudinal control, t is the discrete time and $R_i[\cdot]$ is the reward for using $r(t), j(t)$ for a_i , beginning with the environment and state of the vehicle of a_i .

Eq. (1) emphasizes the choice of instantaneous control r_0, j_0 over the future $r(t), j(t)$. $Q_i(r_0, j_0)$ is a topographic encoding of action values that tells the value of choosing r_0, j_0 for each a_i .

Fig. 4 shows that each inverse model is made up of excitatory and inhibitory loops. The formers activate regions of the r_0, j_0 plane corresponding to free navigable spaces. The latter loops suppress subregions that correspond to close encounters (yellow) or collisions with obstacles (red), (Da Lio et al. (2020), Section III.C).

For action selection, the individual Q_i are aggregated into a single Q function.

$$Q(r_0, j_0) = \max(w_i Q_i(r_0, j_0), i \in \text{affordances}). \quad (2)$$

where w_i are weights that can steer affordance selection. We have shown how this weighting mechanism can be used to navigate the landscape of affordances proactively and implement legal rules in (Da Lio et al. (2020)), long-term prudent behavior (Rosati Papini, Plebe, Da Lio, & Donà, 2022), and to realize human-vehicle interactions (Da Lio, Donà, Rosati Papini, & Plebe, 2022b).

The final action selection is carried out with *multihypothesis sequential probability ratio test algorithm* (MSPRT) (Baum & Veeravalli, 1994; Draglia, Tartakovsky, & Veeravalli, 1999), which is an algorithm that maximizes the probability of making the highest-salience choice in noisy situations and within a maximum decision time constraint.

4. Emergent behaviors

This section demonstrates emergent behaviors. These can be quite complex and look like what one might logically expect. The section presents six examples (simulated in IPG CarMaker and post-processed in Wolfram’s Mathematica): merge, follow and overtake, reaction to cut-in, incorrect pedestrian crossing, traffic not giving way, and navigating an unusual wide lane with parallel traffic and an obstacle. These examples are presented in a supplement video, and the salient moments are commented on here. When possible, we compare the traditional sense-think-act solutions.

4.1. Format of videos

Fig. 5 shows the format of the videos. There are several data visualizations organized in panes. On the top left, the bird’s eye view pane shows the legal corridors (the portion of the road that the self-driving vehicle is legally allowed to use) and other objects on the road. On the top right, there is the corresponding camera view showing the whole road (in this example, the legal corridor is the rightmost lane) and the same objects present in the bird’s eye view. In the upper middle, a panel representing the longitudinal intention/state of the self-driving agent is shown. It shows four alternative conditions that are: (a) free flow, if the vehicle longitudinal control is not limited by objects, legal speed limits, or curves, (b) car-following (when the longitudinal control is constrained by the need to comply with an obstacle ahead, not necessarily a car), (c) speed limit (when the speed must be adapted to the next legal limit), (d) curve (when the speed must be reduced for a next curve). The current state is highlighted. At the bottom left, there is a panel concerned with the lateral state/intention. The panel shows three distinct affordances: remaining on the current legal corridor or moving to an adjacent corridor to the left or right. In the example, there are no legally affordable corridors on both sides; hence, the squares are gray. When there are choices (see Section 4.3 in frame labeled “3”), they are shown here. The chosen affordance is highlighted and the others are dimmed.¹ For every affordable corridor, there are two

sub-panels. The bottom one shows a representation of the 3D salience map (the top one is an intermediate step in the computation of the lateral salience map, which is related to the estimated probability of time-to-lane-crossing (TTLC) given a particular lateral control choice — see Appendix). On the right, there are two final panes. In the middle, the lateral salience, which visualizes the range of lateral control for remaining in the corridor. At the bottom, a chart shows the inhibitions (absolute in red, partial in yellow) produced by the obstacles shown in the bird’s eye and camera views. The inhibited regions indicate the combination of lateral and longitudinal control at risk of collision. The lateral salience map and the inhibition map are related to the selected affordance (highlighted in the lateral state pane). The lateral salience is the lateral cross section of the 3D salience map. The inhibition map is a contour density map of the 3D salience, where red means zero salience. The green-white boundary on the inhibition map visualizes the position of the maximum longitudinal salience.

A tiny blue circle indicates the instantaneous longitudinal and lateral control selected by the agent. This choice, propagated backward in the lateral state and longitudinal state panels, identifies which loop in the subsumption architecture corresponds to the chosen control, i.e., the agent’s longitudinal and lateral intentions (the reactivation of the subsumed loops permits the generation of behavior in every detail).

An interesting observation related to the inhibition chart is that the inhibited neurons (in a neural implementation) corresponding to the obstacles may influence the agent’s decision: if the inhibitions are close enough, they reduce the likelihood of the choice (close actions reinforce each other, and distant actions compete against each other). Between choices with close inhibitions and similar-valued options without inhibitions in the neighborhood, the agent will tend to prefer the latter, achieving robust behavior.

4.2. Emergent merge behavior

Fig. 6 shows three salient moments of the merge video example. In frame “1”, two vehicles on the motorway (not shown yet in the camera view) are going to collide with the self-driving vehicle: they produce inhibitions that approach the action chosen by the agent. In frame “2”, the inhibitions of the obstacles overlap the range of lateral control that permits the vehicle to remain in the lane. The only option for the agent to stay in the lane is to choose a longitudinal control toward the bottom of the chart, which means decelerating. The selected action is also shown in the 3D salience chart. The chosen longitudinal control is at the bottom edge of the partial inhibition, automatically producing an emergent car-following behavior. As the gap opens, the self-driving vehicle resumes speed approaching the speed limit and merges with the motorway (“3”). Similar examples found in the literature may be (Kreutz and Eggert (2021), Menéndez-Romero (2021)). They use an elaborate algorithm to produce an analogous behavior.

4.3. Follow and overtake

Fig. 7 shows three salient frames of the follow and overtake example. At “1”, after completing the approach phase, the self-driving vehicle enters the car-following state (the selected action touches the obstacle inhibition and it would not be possible to travel faster). The two vehicles ahead proceed at the same speed (longitudinal inhibitions have a matching bottom), and there is no added value in moving to the left lane. The change left affordance exists but is not selected (dimmed colors). At “2”, the obstacle on the left accelerates. Immediately, its inhibition in the control space moves upward and the change left affordance wins (the winner peak is shown in the lateral state panes). Therefore, the intention of the agent changes to the left lane change (the planned trajectory is shown in the bird’s eye view), while the longitudinal state remains “car following” (following the left car). When the lane change is completed at “3”, the agent has two options: staying on the left lane or returning to the right. The latter wins (lanes

¹ There is a fourth affordance not shown in the pane, which corresponds to the entire drivable road. It represents a nonlegal option that has very low priority and is used only when legal rules must be broken to avoid accidents (see Da Lio et al. (2020), Section III.C)).

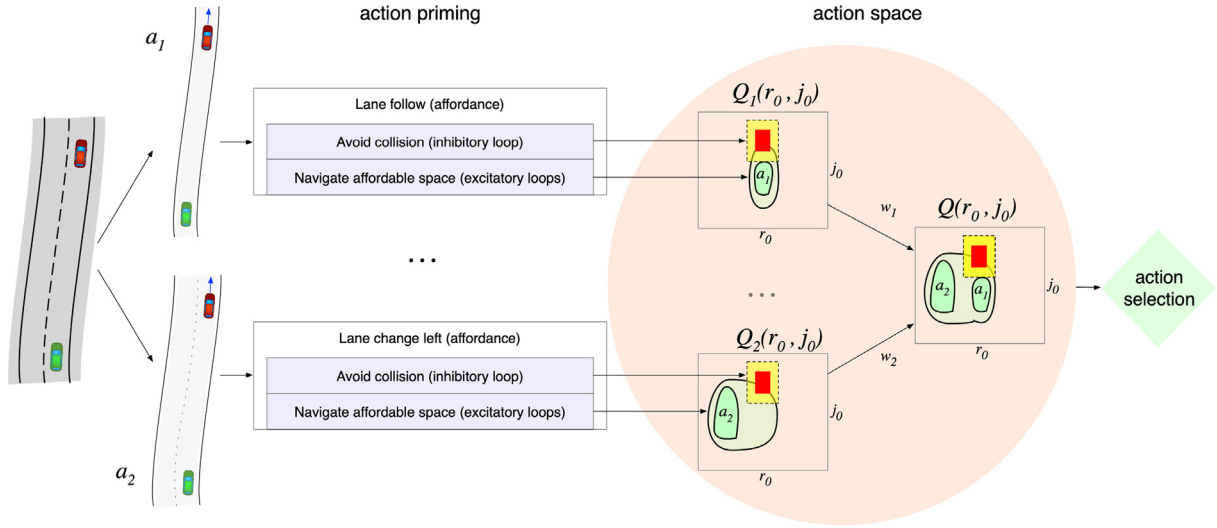


Fig. 4. Dreams4Cars implementation of layered control with affordance competition. In this example, the green car has two options: traveling in the lane (a_1) or moving to the left lane (a_2). For each option, a map is created, which represents the reward $Q_i(r_0, j_0)$ for using the instantaneous control r_0, j_0 (respectively, lateral and longitudinal control). The contour lines on the reward map are filled with green. The more intense the color, the higher the value, which peaks around the optimal control choice. Red and yellow represent regions inhibited by obstacles (with zero or diminished value). The value maps are aggregated into a final decision map $Q(r_0, j_0)$, weighted by weights w_i that can be used for a variety of purposes described in the text and [Appendix A](#).

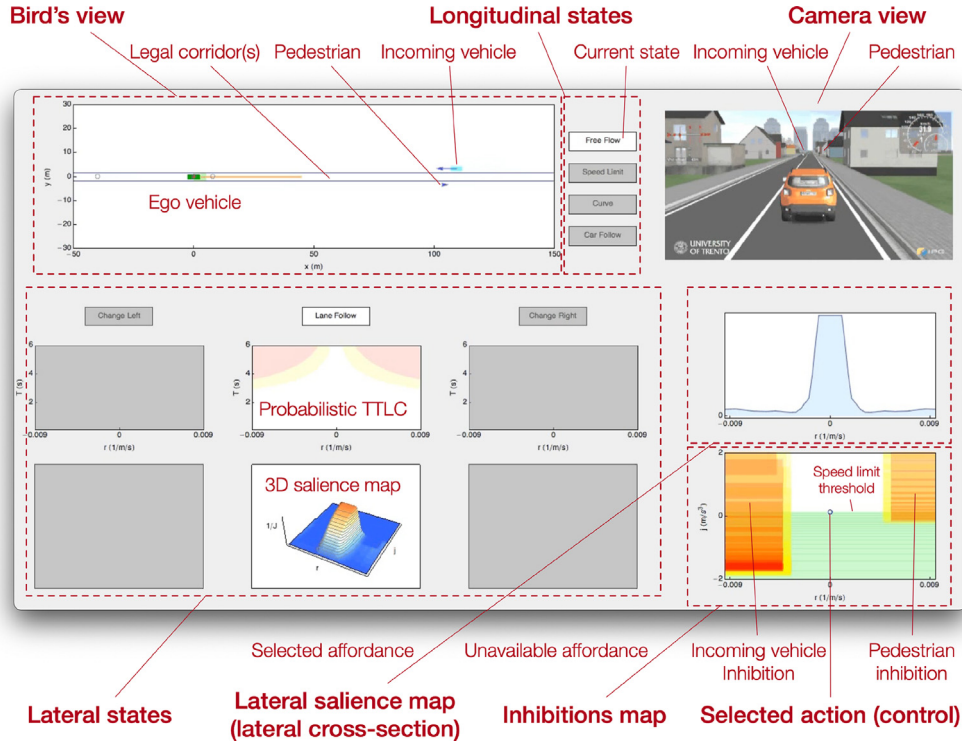


Fig. 5. Format of videos. This graphical interface summarizes the agent's computations using several real-time data charts. The bird's view shows a scheme of the relevant objects in the environment. The axes labels are: "x (m)" forward direction in meters (range -50 m to 150 m), and "y (m)" lateral direction in meters (range -30 m to $+30$ m). The longitudinal states box shows four conditions that affect the speed choice: free flow, speed limit, curve, and car-following. The inhibition map is the aggregated salience chart $Q(r_0, j_0)$ of [Fig. 4](#). The axes labels are: "r (1/m/s)" lateral control in the trajectory curvature rate (m^{-1}/s) (range -0.009 m^{-1}/s to 0.009 m^{-1}/s) and "j m/s^3 " longitudinal — jerk — control (range -2 m/s^3 to $+2$ m/s^3). The lateral states box shows the $Q_i(r_0, j_0)$ for the three legal lanes. The probabilistic TTLC (Time To Lane Crossing) subgraph shows the time when the probability of leaving the corridor in absence of corrections exceeds two thresholds (see [Appendix](#)). The axes labels are: "r (1/m/s)" and "T (s)" (range 0 to 6 s). The 3D salience map shows $Q_i(r_0, j_0)$ with the elevation axis (labeled "1/J") being the salience (see [Appendix](#)). Finally, the lateral salience map shows the function $f(r_0)$ (see [Appendix](#)).

to the right are slightly biased to favor the legal recommendation to return to the right when possible). The winner peak in the lateral intention pane corresponds to a path that will reenter the right lane

after clearing the obstacle to the right, as can be seen in the video. Note that the intention of returning to the right after completing the overtake is taken very early, well before the right obstacle is cleared.

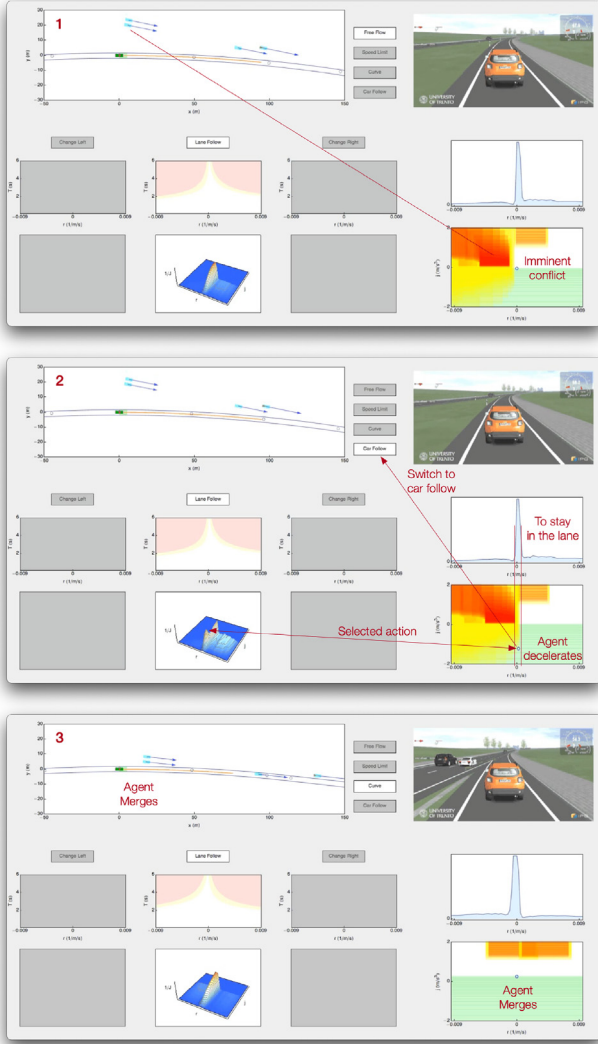


Fig. 6. Emergent negotiation in a merge scenario.

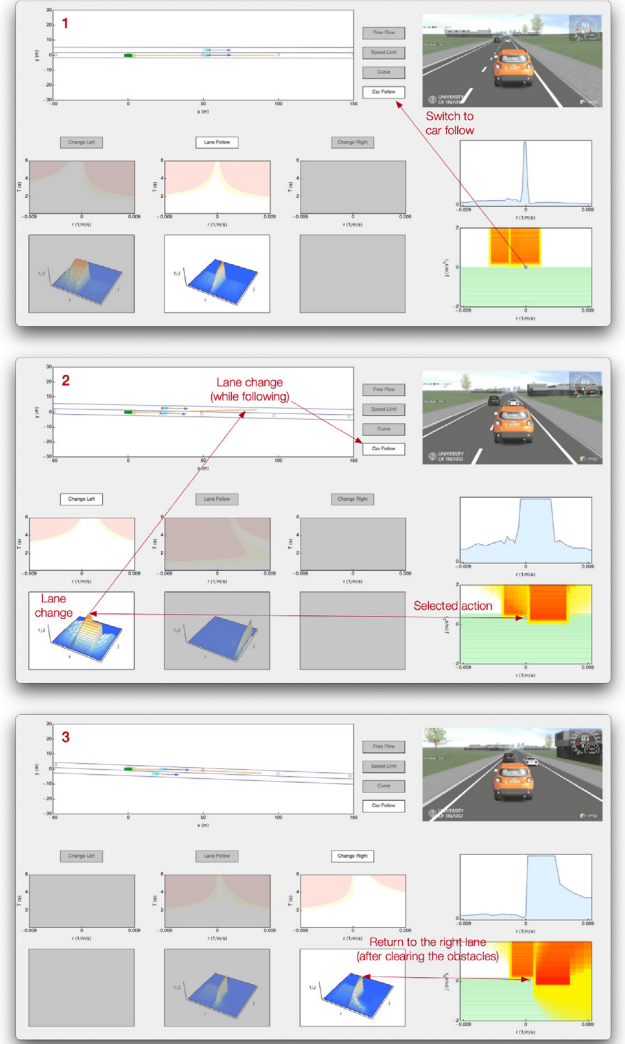


Fig. 7. Alternating opportunistic overtaking and car-following.

4.4. Reacting to cut-in

Fig. 8 illustrates the reaction to a cut-in maneuver. At “1”, a vehicle arriving from the rear in the fast left lane is detected first. It blocks the lane change option (only the trajectories that do not complete the lane change survive on the salience map of the left panel.²) At “2”, the fast vehicle frees the lane: the salience map for the left lane change widens, including complete lane change maneuvers. However, the agent does not move leftwards, as the current lane is free. At “3”, the obstacle decelerates and moves toward our lane (cut-in). The winner action is now the lane change: at “3”, as soon as the predicted obstacle path indicates cut-in, the agent plans to overtake. The rest of the video shows the completion of the overtake maneuver with a return to the right lane, similar to the example in Section 4.3.

4.5. Pedestrian crossing incorrectly

Corner cases may arise from the incorrect behavior of other road users (it is difficult to imagine every possible incorrect behavior).

² Recall that left and right affordances include both the current and the nearby lane.

In this example, a pedestrian crosses the road when the traffic light turns red for the pedestrian and green for the self-driving vehicle. The salient frames are shown in Fig. 9. At “1”, the self-driving vehicle stops because of the inhibition produced by the red traffic light. Crossing traffic produces additional inhibition to the same motor space region. When the traffic light turns green (“2”), and the crossing traffic stops, that motor region would become free, allowing the vehicle to start. However, the movement of the pedestrian keeps that region re-inhibited (in fact, the path of a moving entity is predicted regardless of whether it is correct, as explained in the Appendix). Thus, the self-driving vehicle does not start regardless of the green traffic light. We had a chance to compare the same situation with traditional self-driving software. The comparison is shown at EU project Dreams4Cars (2018). The traditional system starts when the traffic light turns green and brakes when the pedestrian is about to enter the lane. This example does not mean that every traditional self-driving system will behave the same way in that situation. However, it indicates that rule-based systems may be fragile and rigid.

4.6. Traffic not giving way

Fig. 10 shows a similar case: the crossing traffic does not give way. A safe stop maneuver emerges from the inhibition caused by the crossing vehicles.

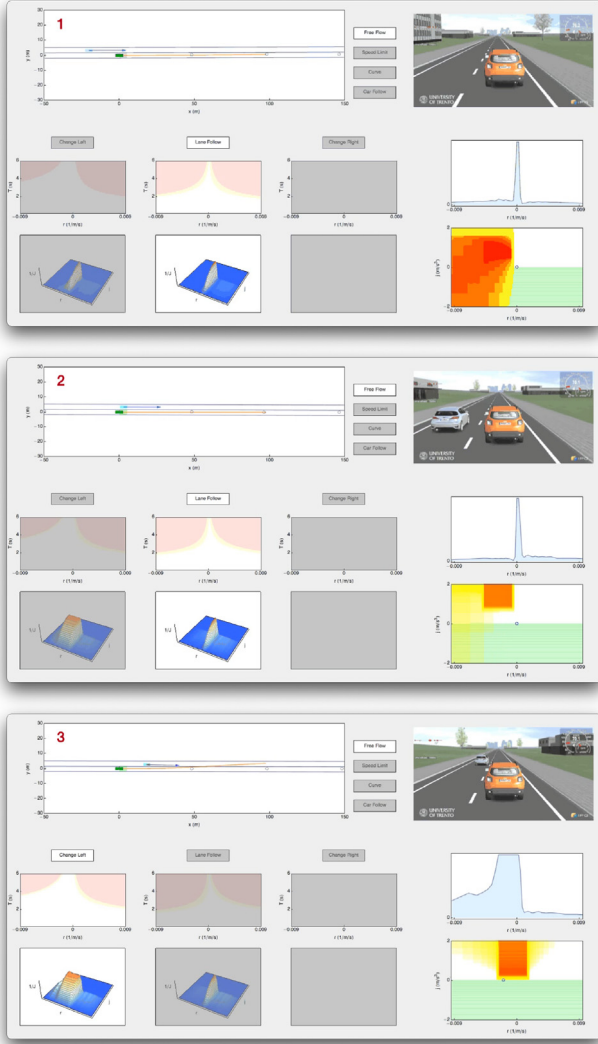


Fig. 8. Emergent reaction to cut-in.

4.7. Unusual wide lane with parallel traffic and obstacle

The situation shown in Fig. 11 combines a few unusual elements. The lane is 25 m wide. There is a stationary vehicle far ahead in the lane (with an affordance to room to pass on either side). Two vehicles follow the self-driving vehicle closely on both sides, at slightly different distances. The traditional self-driving software used for the pedestrian example (Section 4.5) would stop the car behind the stopped vehicle because it is a busy lane (without considering the unusual width of the lane). Such behavior is dangerous because people would not expect that (we again remark that this is the case of the particular implementation that we could use for Section 4.5 and here). Common sense would suggest letting the closest rear car on the left pass and then moving on the left to pass the stopped obstacle on the left. This behavior is emergent, and explained with Fig. 11 and Fig. 12. Vehicles start from rest. At low speed (time 5 s in Fig. 12) the entire action space is affordable. The agent accelerates under free flow conditions, remaining in the center of the lane (the vehicle speed is visible on the tachometer in the camera view). At about 8.4 s, the nearby vehicles — with increased speed — restrict the affordable control choices. At about 8.5 s the front obstacle is detected, which divides the affordable space into two parts (a_1 and a_2). Between 8.5 s and 8.8 s (frame “1” in Fig. 11 is between the two times) the agent opts for the faster a_1 ,



Fig. 9. Pedestrian crossing the road illegally.

Fig. 10. Stop in traffic that does not yield.

corresponding to passing in front of the right car and to the right of the obstacle. However, the vehicle on the right is getting closer until, at time 9.6 s, a_1 vanishes. At about time 10 s the agent reverts the decision to a_1 , which is tailing (and stopping) behind the front obstacle (frame “2”). As the speed decreases, at time 12 s the left vehicle is fast enough to create another affordance, a_3 (frame “3”), which corresponds to tailing the left vehicle and passing to the left of the obstacle (note that the left vehicle is still on the side, but the motor space is based on future predicted position and therefore anticipates possible actions).

5. Discussion and conclusion

Section 3 presented a self-driving agent characterized by creating emergent behaviors, and we gave six examples of emergent behaviors (from Section 4.2 to Section 4.7). The principles used here are borrowed from biology and can be seen as a proposed step toward AGI, an alternative to the traditional narrow AI approaches mentioned in Section 2.

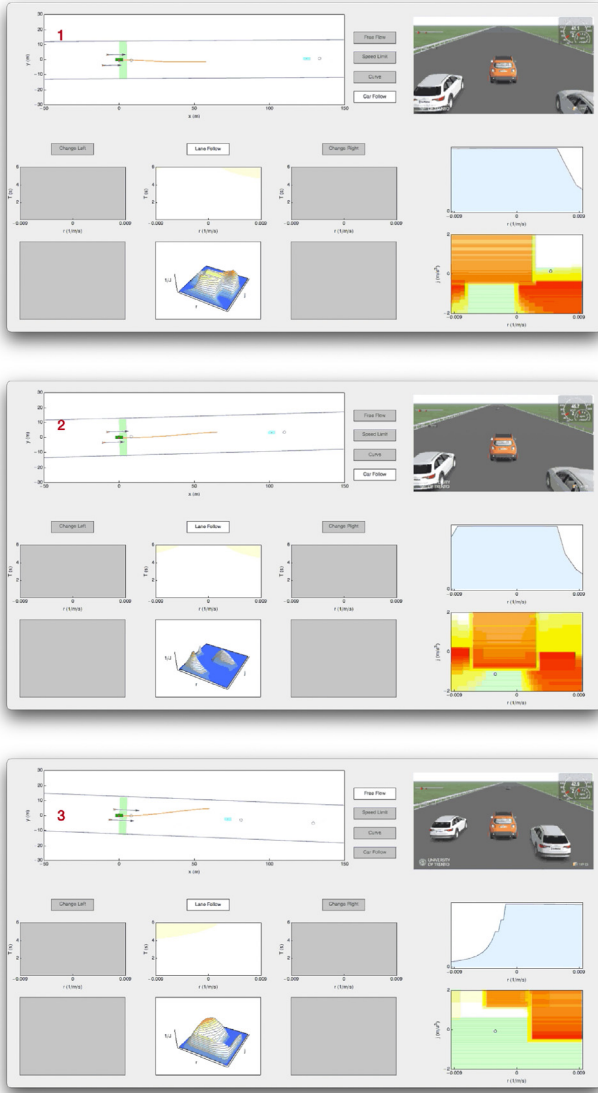


Fig. 11. Complex emergent behavior.

The implementation of the agent as described in Section 3.1, is limited to the strict extent necessary for this paper; but a complete description is found in Da Lio et al. (2020). Here, we want to comment on two points:

- (1) The biological principles exploited in the agent and their implications for the agent's operation.
- (2) How this approach compares to the ones based on narrow AI; and what might be the advantages for development/maintenance and scalability.

5.1. Biological principles exploited in the agent

Cisek's Affordance Competition hypothesis (Cisek, 2007) is the main idea. Parallel action priming and selection through a competitive process is the theoretical element that informs the agent's sensorimotor system (Fig. 4 vs. Cisek (2007, Fig. 1)). Continuous competition between affordances forms the basis for emergent adaptive behavior.

Another important idea is the topographic organization of the motor space, where affordances are encoded by hills of activity with the heights of the hills that encode the action value (or the salience, (2)).

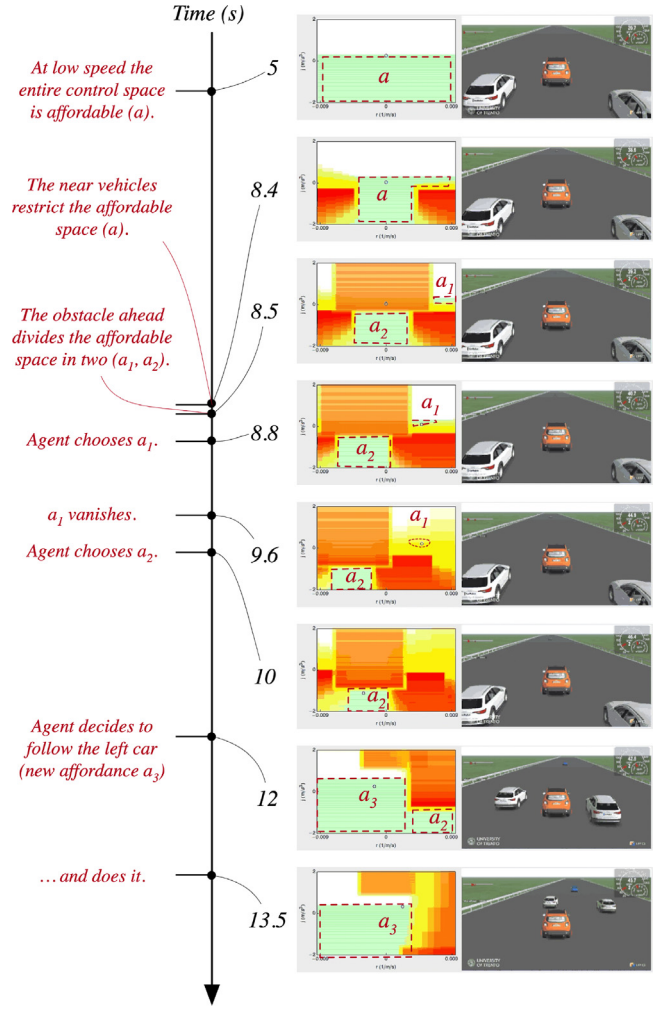


Fig. 12. Timeline of events and decisions..

Trajectories that begin with similar control lie close to each other in the motor space and reinforce each other. Hence, by selecting one particular instantaneous control, the agent commits to a family of future maneuvers (beginning with a similar control and possibly diverging later) rather than selecting just one. This minimum commitment principle further maximizes adaptability because alternative maneuvers with high values may exist close to the highest-valued one.

Finally, the whole system is organized hierarchically (Fig. 4) with different branches dedicated to high-level goals. The branches are aggregated into a unique action-value map $Q(r_0, j_0)$ with weights that allow one to navigate the affordance landscape according to Pezzulo's hypothesis (Pezzulo & Cisek, 2016). This allows the agent's behavior to be guided by high-level directives, effectively modulating the value function $Q(r_0, j_0)$. This flexibility (i.e., a modifiable reward function) does not normally exist in traditional reinforcement learning implementations. Even if not shown here, we must mention that this mechanism was used for the implementation of traffic regulations (Da Lio et al., 2020), long-term cautious behavior (Rosati Papini et al., 2022), and human-agent interactions (Da Lio et al., 2022b). Thus, the emergent behavior characteristic of this agent does not prevent controlling the agent's higher-level behaviors with directives (or rules), for example, to comply with a highway code.

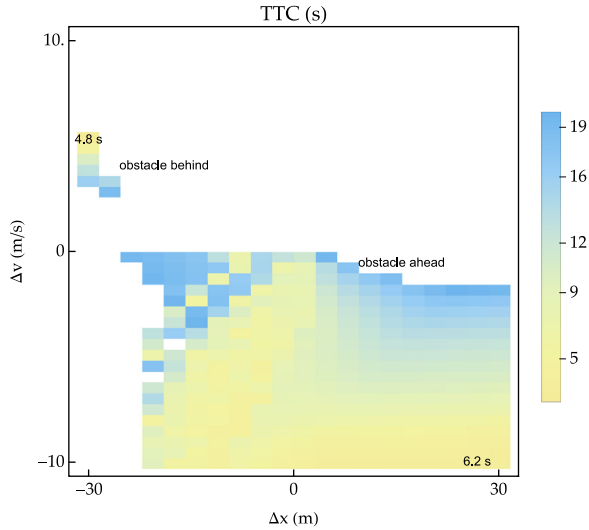


Fig. 13. Speed adaptation on ramps obtained as an emerging behavior: time to collision (TTC) after merging with different relative longitudinal positions (Δx) and different relative speed (Δv). The graph is directly comparable to (Kreutz & Eggert, 2021, Fig. 7).

5.2. How does this approach compare to narrow AI ones

Most traditional motion planning methods work by selecting one trajectory from a pool of alternatives; however, not with the sophistication of the above principles. In particular, while local trajectory optimization is obtained with various methods, among which optimal control, the choice among alternative higher-level behaviors is typically produced with ad hoc algorithms.

As an example, consider the merge scenario Section 4.2. An algorithm to deal with this situation was mentioned earlier (Kreutz & Eggert, 2021). It adapts the Intelligent Driver Model (IDM) to modulate the speed of a vehicle on a ramp entering the main road. The paper proposes a modified IDM (GIDM) and shows, with various performance indicators, that it works to regulate vehicle entrance. Although the ramp is adjacent to the road, the vehicle enters the road only at the end; therefore, it is equivalent to the situation in Section 4.2. We replicated the study with our agent, performing simulations spanning the same conditions described in the paper (Kreutz & Eggert, 2021, Section VI.B) to find that a fine adaptation of the speed in the ramp is automatically obtained for all the conditions considered. For example, Fig. 13 shows the time-to-collision (TTC) after merge obtained with different relative longitudinal positions of the vehicle and obstacle (Δx) and with different relative speeds (Δv). The graph is directly comparable to (Kreutz & Eggert, 2021, Fig. 7). The top-left region occurs when the agent merges in front of the obstacle. The bottom-right region occurs when the agent merges behind (this latter part is not considered in Kreutz and Eggert (2021)). The white color occurs when the tailing vehicle (either the agent or the obstacle) is slower than the leading vehicle (therefore, TTC does not exist), or TTC is greater than 20 s.

This example shows that the principles commented on in Section 5.1 may be an alternative to programmed algorithms. Algorithms are based on schematic situations and may have limitations. For example, as discussed in Kreutz and Eggert (2021, Section IV) the algorithm is an adaptation of one-dimensional dynamics to the reality and needs various schematizations which are explained in subsections B to E. Of course, more complex algorithms are possible, for example (Menéndez-Romero, 2021), which are more difficult to develop, test, and validate. Thus, in our view, having behaviors that emerge from biological principles is an attractive perspective.

5.3. Limitations of this study and future plans

The main limitation of this study is its anecdotal nature. On the one hand, we show nontrivial emergent behaviors. On the other hand, the limited number of examples given here does not permit one to claim that emergent behaviors will be correct for any possible situation (but this problem holds also for algorithms). The testing and validation of self-driving software for long-term driving is a major undertaking for current autonomous driving development programs. For example, the recent EU project SUNRISE (Safety as a URaNce fRamework for connected, automated mobility SystEms) (EU project SUNRISE, 2022) aims at developing test scenarios and test cases to address validation in the long term, inheriting the findings of many previous projects.

Note that, as explained in the Introduction, the presented work proposes a bridge between narrow AI and AGI, although it is not possible for it to fully belong to the latter. This is because the agent must give importance to the restricted set of affordances that are meaningful in the driving context. On the other hand, in the case of a human driver, behaviors can emerge from the knowledge of a wider variety of contexts, in which the social aspect has a significant weight. Despite this limitation, the agent architecture shown here produces behaviors that would otherwise require careful programming (Section 5.2, and Sections 4.5, 4.7 where they mention the outcome of a traditional software).

So what if we discover situations where the behaviors are not acceptable? Algorithmic solutions would call for modification of the algorithms and re-testing. In contrast, the agent above can exploit the affordance navigation and biasing mechanism (Section 5.1) to learn better high-level choices. An example of how this may work was given in Rosati Papini et al. (2022) for learning the speed choice in the case of distracted pedestrians.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data availability

The data that has been used is confidential.

Appendix A

This appendix provides additional information on our implementation (to be considered as an example) and links to previous literature.

The agent recognizes up to three legal corridors (Da Lio et al., 2020, Section II-C.1,3): staying in the current lane or moving to a nearby lane on the left or right (we assume that a double lane change is always carried out in two single lane changes). In addition to the legal corridors, the whole drivable surface is also considered a fourth possibility, but the weight (w in Fig. 4) assigned to this option is very low, to use it only as a last resource if there are no collision-free solutions in the legal lanes.

For each of the (up to) four affordable corridors, the agent evaluates the salience $Q_i(r_0, j_0)$ according to principle (1). In practice, we simplify the computation assuming that lateral and longitudinal dynamics are weakly coupled, which is a usual simplification in non-extreme vehicle dynamics conditions. Hence:

$$Q(r_0, j_0) \cong f(r_0)g(j_0). \quad (3)$$

where $f(r_0)$ is the lateral salience and $g(j_0)$ the longitudinal salience.

The lateral salience is evaluated as follows: for every instantaneous control, r_0 , a stochastic motion model,³ is considered to produce a family of possible trajectories, some of which may trespass the lane borders. The lateral salience is taken as the time when the probability of leaving the lane exceeds a low threshold: the longer this is, the more time the agent has to carry out corrections⁴ and the easier and smoother the lane-following task is. For $f(r_0)$ we rescale the TTLC (and in particular subtract 2 s from the time).

The longitudinal salience is evaluated as follows: a simplified version of the human longitudinal motion model (calibrated on experimental data) (Da Lio, Mazzalai, Gurney, & Saroldi, 2018b) is considered and, for every instantaneous longitudinal control j_0 , various subgoals are examined (e.g., stopping at some stop line; see examples in Da Lio, Mazzalai, and Darin (2018a)). The cost (Da Lio et al., 2018b, Equation 4) is used for the inverse of $g(j_0)$. Obstacles are managed in two steps (Da Lio et al., 2020, Section II-C.5): first prediction of their likely path and then inhibition of the motor space regions leading to the spatio-temporal locations (for the latter step, motor primitives for the lateral and longitudinal control above are reused). For standing pedestrians (who do not have a predicted path), a cautious approaching speed is obtained by reinforcement learning considering how they might move (Rosati Papini et al., 2022).

For decision-making, the value maps of the individual corridors are aggregated into a unique action-value motor chart with (2). In the aggregation process, the original rewards Q_i are weighted with weights w_i . These weights may be scalar quantities (in which case decision making is steered toward one of the i th affordances, or they may be functions $w_i(r_0, j_0)$ that can be used to bias particular regions on the motor space (for example, faster travel). The idea of steering decision-making in this way derives from Pezzulo and Cisek (2016). Biasing is intrinsically safe because collision trajectories are completely inhibited in Q_i and cannot be selected regardless of the biases w_i . A discussion of this mechanism is found in Da Lio et al. (2020, Section II-E), while Da Lio et al. (2020, Section III-C) demonstrates using proactive lane change on motorways and Da Lio et al. (2022b) demonstrates emergent human-robot interactions based on the same principle.

Appendix B. Supplementary data

Supplementary material related to this article can be found online at <https://doi.org/10.1016/j.cogsys.2022.12.007>.

References

- Adams, E. (2020). Why we're still years away from having self-driving cars. *Vox*, <https://www.vox.com/recode/2020/9/25/21456421/why-self-driving-cars-autonomous-still-years-away>.
- Anderson, M. (2020). Surprise! 2020 is not the year for self-driving cars. *IEEE Spectrum*, <https://spectrum.ieee.org/transportation/self-driving/surprise-2020-is-not-the-year-for-selfdriving-cars>.
- Arieff, A. (2021). The underwhelming reality of driverless cars. *City Monitor*, <https://citymonitor.ai/transport/the-underwhelming-reality-of-driverless-cars>.
- Bansal, M., & Ogale, A. (2018). Learning to drive: Beyond pure imitation. *Waymo Research, Medium*, <https://medium.com/waymo/learning-to-drive-beyond-pure-imitation-465499f8bcb2>.
- Barsalou, L. W. (2008). Grounded cognition. *Annual Review of Psychology*, 59, 617–645.
- Baum, C., & Veeravalli, V. (1994). A sequential procedure for multihypothesis testing. *IEEE Transactions on Information Theory*, 40(6), 1994–2007.
- Boggs, A. M., Arvin, R., & Khattak, A. J. (2020). Exploring the who, what, when, where, and why of automated vehicle disengagements. *Accident Analysis and Prevention*, 136, 105406.
- Bojarski, M., Del Testa, D., Dworakowski, D., Firner, B., Flepp, B., Goyal, P., et al. (2016). End to end learning for self-driving cars. *ArXiv preprint arXiv:1604.07316*.
- Bojarski, M., Yeres, P., Choromanska, A., Choromanski, K., Firner, B., Jackel, L., et al. (2017). Explaining how a deep neural network trained with end-to-end learning steers a car. *ArXiv preprint arXiv:1704.07911*.
- Brooks, R. (1986). A robust layered control system for a mobile robot. *IEEE Journal of Robotics and Automation*, 2(1), 14–23.
- Brooks, R. A. (1991). Intelligence without representation. In D. Kirsh (Ed.), *Artificial Intelligence*, 47(1811), 139–159.
- CB Insights (2015–2020). 40+ Corporations working on autonomous vehicles. In *CB insights research briefs*. <https://www.cbinsights.com/research/autonomous-driverless-vehicles-corporations-list/>.
- Cisek, P. (2007). Cortical mechanisms of action selection: the affordance competition hypothesis. *Philosophical Transactions of the Royal Society, Series B (Biological Sciences)*, 362(1485), 1585–1599.
- Claussmann, L., Revilloud, M., Gruyer, D., & Glaser, S. (2019). A review of motion planning for highway autonomous driving. *IEEE Transactions on Intelligent Transportation Systems*, 1–23.
- Connected Automated Driving (2021). *Workshop on Edge Cases for Automated Driving*, <https://www.connectedautomateddriving.eu/blog/event/workshop-on-edge-cases-for-automated-driving/>.
- Da Lio, M., Donà, R., Rosati Papini, G. P., Biral, F., & Svensson, H. (2020a). A mental simulation approach for learning neural-network predictive control (in self-driving cars). *IEEE Access*, 8, 192041–192064.
- Da Lio, M., Donà, R., Rosati Papini, G. P., & Gurney, K. (2020). Agent architecture for adaptive behaviours in autonomous driving. *IEEE Access*, 8, 154906–154923.
- Da Lio, M., Donà, R., Rosati Papini, G. P., & Plebe, A. (2022b). The biasing of action selection produces emergent human-robot interactions in autonomous driving. *IEEE Robotics and Automation Letters*, 7(2), 1254–1261.
- Da Lio, M., Mazzalai, A., & Darin, M. (2018a). Cooperative intersection support system based on mirroring mechanisms enacted by bio-inspired layered control architecture. *IEEE Transactions on Intelligent Transportation Systems*, 19(5), 1415–1429.
- Da Lio, M., Mazzalai, A., Gurney, K., & Saroldi, A. (2018b). Biologically guided driver modeling: the stop behavior of human car drivers. *IEEE Transactions on Intelligent Transportation Systems*, 19(8), 2454–2469.
- De Haan, P., Jayaraman, D., & Levine, S. (2019). Causal confusion in imitation learning. *arXiv:1905.11979 [cs, stat]*.
- Dixit, V. V., Chand, S., & Nair, D. J. (2016). Autonomous vehicles: Disengagements, accidents and reaction times. *PLoS One*, 11(12), e0168054.
- Draglia, V., Tartakovsky, A., & Veeravalli, V. (1999). Multihypothesis sequential probability ratio tests .I. Asymptotic optimality. *IEEE Transactions on Information Theory*, 45(7), 2448–2461.
- Eliot, L. (2021). Whether those endless edge or corner cases are the long-tail doom for AI self-driving cars. *Forbes, Section: Transportation*, <https://www.forbes.com/sites/lanceeliot/2021/07/13/whether-those-endless-edge-or-corner-cases-are-the-long-tail-doom-for-ai-self-driving-cars/>.
- EU project COSPAL (2007). Cognitive systems using perception-action learning. *FP6*, <https://cordis.europa.eu/project/id/004176>.
- EU project DIPLECS (2010). Dynamic interactive perception-action learning in cognitive systems. *FP7*, <https://cordis.europa.eu/project/id/215078>.
- EU project Dreams4Cars (2018). Dreams4cars comparison with sense-think-act agent. <https://www.youtube.com/watch?v=kONEmwIB1gU>.
- EU project Dreams4Cars (2019). Dream-like simulation abilities for automated cars. H2020 grant agreement no. 731593, <https://www.dreams4cars.eu/en>.
- EU project HEADSTART (2019). Harmonised European solutions for testing automated road transport. EU Grant Agreement No. 824309 <https://www.headstart-project.eu/about/>.
- EU project SUNRISE (2022). Safety assurance framework for connected, automated mobility systems. <https://cordis.europa.eu/project/id/101069573>.
- Gibson, J. (1986). The theory of affordances. In *The ecological approach to visual perception* (pp. 127–143). Lawrence Erlbaum Associates.
- Goertzel, B. (2014). Artificial general intelligence: Concept, state of the art, and future prospects. *Journal of Artificial General Intelligence*, 5, 1–46.
- Goertzel, B., Iklé, M., & Wigmore, J. (2012). The architecture of human-like general intelligence. In P. Wang, & B. Goertzel (Eds.), *Theoretical foundations of artificial general intelligence* (pp. 123–144). Amsterdam: Atlantis Press.
- Goertzel, B., & Pennachin, C. (2006). *Artificial general intelligence*. Berlin: Springer-Verlag.
- González, D., Pérez, J., Milanés, V., & Nashashibi, F. (2016). A review of motion planning techniques for automated vehicles. *IEEE Transactions on Intelligent Transportation Systems*, 17(4), 1135–1145.
- Kalra, N., & Paddock, S. M. (2016). Driving to safety: How many miles of driving would it take to demonstrate autonomous vehicle reliability? *Transportation Research Part A: Policy and Practice*, 94, 182–193.
- Katrakazas, C., Qudus, M., Chen, W.-H., & Deka, L. (2015). Real-time motion planning methods for autonomous on-road driving: State-of-the-art and future research directions. *Transportation Research Part C (Emerging Technologies)*, 60, 416–442.
- Keijzer, F. (2002). Representation in dynamical and embodied cognition. *Cognitive Systems Research*, 3(3), 275–288.

³ In our case, the stochastic motion model is made of minimum square jerk trajectories with stochastic coefficients. We use a piecewise constant approximation of the corridor curvature for which the minimum lateral jerk trajectories are piecewise quintic polynomials.

⁴ The notion is similar to the concept of time-to-lane-crossing commonly used to evaluate lane keeping.

- Kreutz, K., & Eggert, J. (2021). Analysis of the generalized intelligent driver model (GIDM) for merging situations. In *2021 IEEE intelligent vehicles symposium (IV)* (pp. 34–41).
- lv, C., Cao, D., Zhao, Y., Auger, D. J., Sullman, M., Wang, H., et al. (2018). Analysis of autopilot disengagements occurring during autonomous vehicle testing. *IEEE/CAA Journal of Automatica Sinica*, 5(1), 58–68.
- Marshall, A., & Davies, A. (2019). Uber's self-driving car didn't know pedestrians could Jaywalk. *Wired*, <https://www.wired.com/story/ubers-self-driving-car-didnt-know-pedestrians-could-jaywalk/>.
- Marti, G., Morice, A. H. P., & Montagne, G. (2015). Drivers' decision-making when attempting to cross an intersection results from choice between affordances. *Frontiers in Human Neuroscience*, 8.
- Menéndez-Romero, C. (2021). Maneuver planning for highly automated vehicles. *Albert-Ludwigs-Universität Freiburg*, <https://Freidok.Uni-Freiburg.de/data/222337>.
- Paden, B., Čáp, M., Yong, S. Z., Yershov, D., & Frazzoli, E. (2016). A survey of motion planning and control techniques for self-driving urban vehicles. *IEEE Transactions on Intelligent Vehicles*, 1(1), 33–55.
- Pendleton, S., Andersen, H., Du, X., Shen, X., Meghjani, M., Eng, Y., et al. (2017). Perception, planning, control, and coordination for autonomous vehicles. *Machines*, 5(1), 6.
- Pezzulo, G., & Cisek, P. (2016). Navigating the affordance landscape: Feedback control as a process model of behavior and cognition. *Trends in Cognitive Sciences*, 20(6), 414–424.
- Porrill, J., Dean, P., & Anderson, S. R. (2012). Adaptive filters and internal models: Multilevel description of cerebellar function. *Neural Networks*, 47, 134–149.
- Prescott, T. J., Redgrave, P., & Gurney, K. (1999). Layered control architectures in robots and vertebrates. *Adaptive Behavior*, 7(1), 99–127.
- Project PEGASUS (2019). Project for the Establishment of Generally Accepted quality criteria, tools and methods as well as Scenarios and Situations for the release of highly-automated driving functions. Funded by the German ministry for economic affairs. <https://www.pegasusprojekt.de/en/about-PEGASUS>.
- Rosati Papini, G. P., Plebe, A., Da Lio, M., & Donà, R. (2022). A reinforcement learning approach for enacting cautious behaviours in autonomous driving system: Safe speed choice in the interaction with distracted pedestrians. *IEEE Transactions on Intelligent Transportation Systems*, 7, 8805–8822.
- Seitz, J. A. (2000). The bodily basis of thought. *New Ideas in Psychology*, 18(1), 23–40.
- The New York Times (2021a). Silicon valley is resetting expectations for self-driving cars and settling in for years of more work. *The New York times*, <https://www.nytimes.com/2021/05/25/business/silicon-valley-is-resetting-expectations-for-self-driving-cars-and-settling-in-for-years-of-more-work.html>.
- The New York Times (2021b). The costly pursuit of self-driving cars continues on and on and on. *The New York Times*, <https://www.nytimes.com/2021/05/24/technology/self-driving-cars-wait.html>.
- US National Transport Safety Board (2019). Collision between vehicle controlled by developmental automated driving system and pedestrian. <https://www.nts.gov/news/events/Pages/2019-HWY18MH010-BMG.aspx>.
- Van Gelder, T. (1995). What might cognition be, if not computation? In W. G. Lycan (Ed.), *Journal of Philosophy*, 92(7), 345–381.