

Subjective and objective descriptions of driving scenes in support of driver-automation interactions

Cabrall, Christopher; Happee, Riender; de Winter, Joost

Publication date

2016

Document Version

Final published version

Citation (APA)

Cabrall, C., Happee, R., & de Winter, J. (2016). *Subjective and objective descriptions of driving scenes in support of driver-automation interactions*. Poster session presented at HFES 2016: Annual Meeting Human Factors and Ergonomics Society , Prague, Czech Republic.

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Subjective and Objective Descriptions of Driving Scenes in Support of Driver-Automation Interactions

Christopher D. D. Cabrall, Riender Happee, Joost C. F. de Winter
Delft University of Technology

Introduction

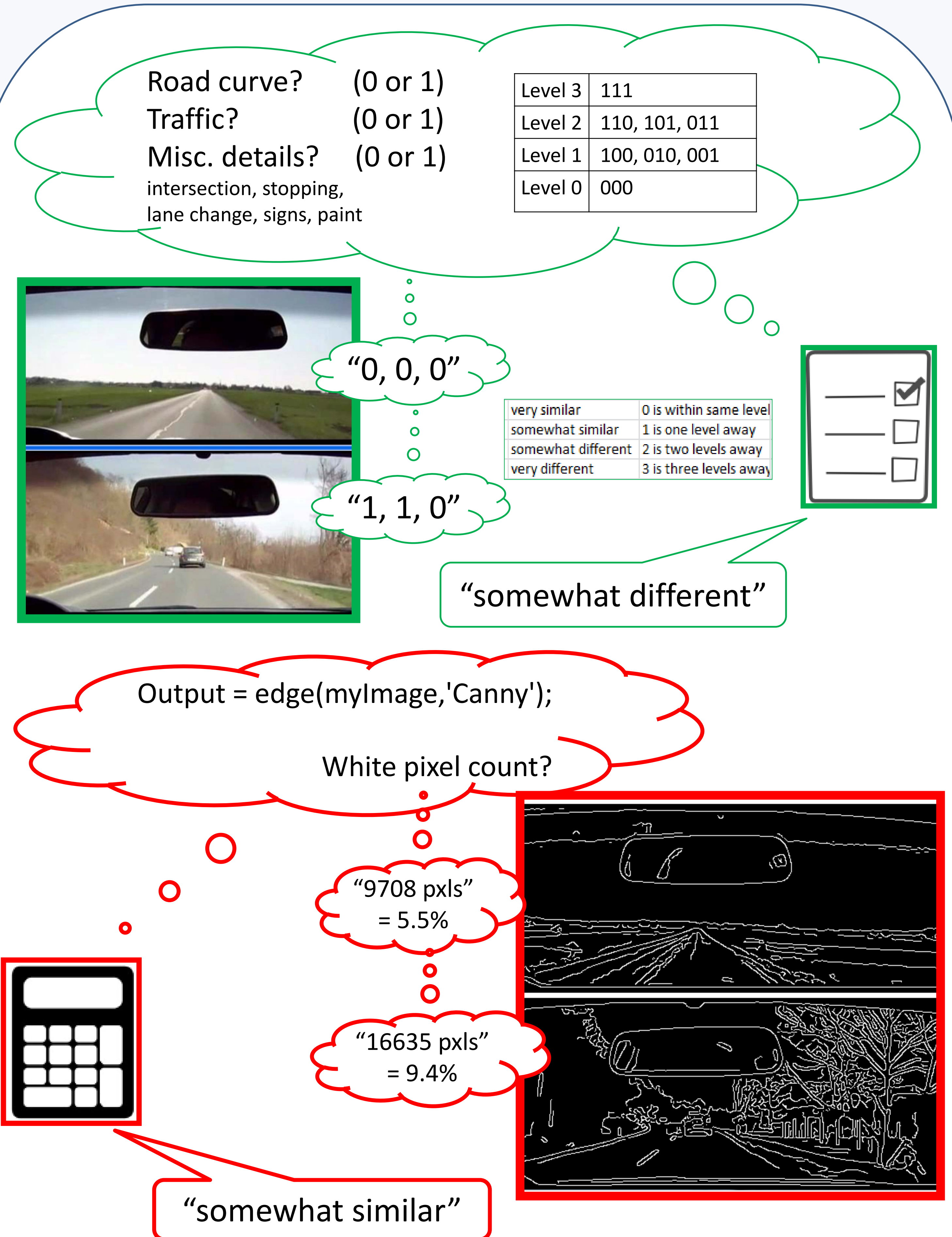
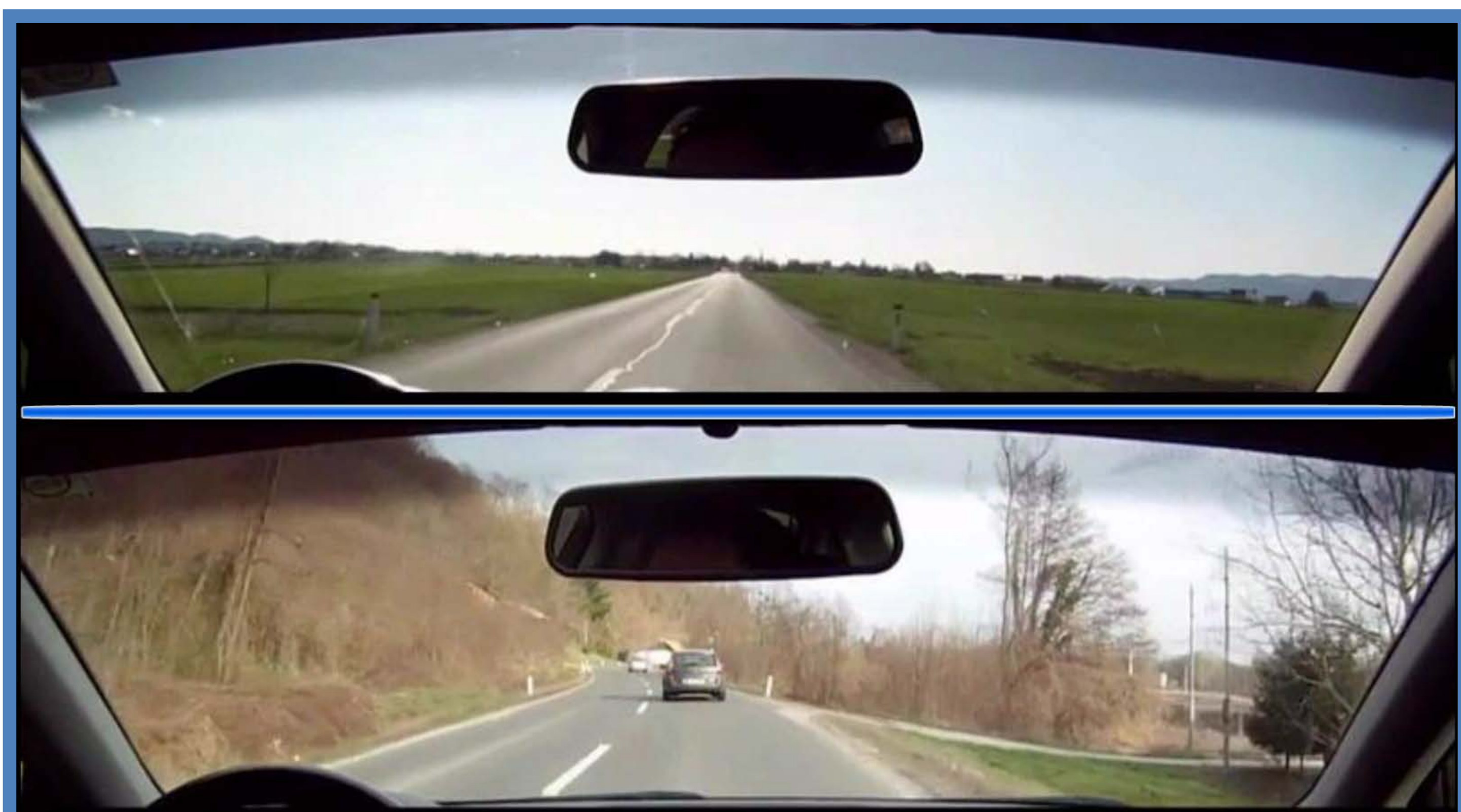
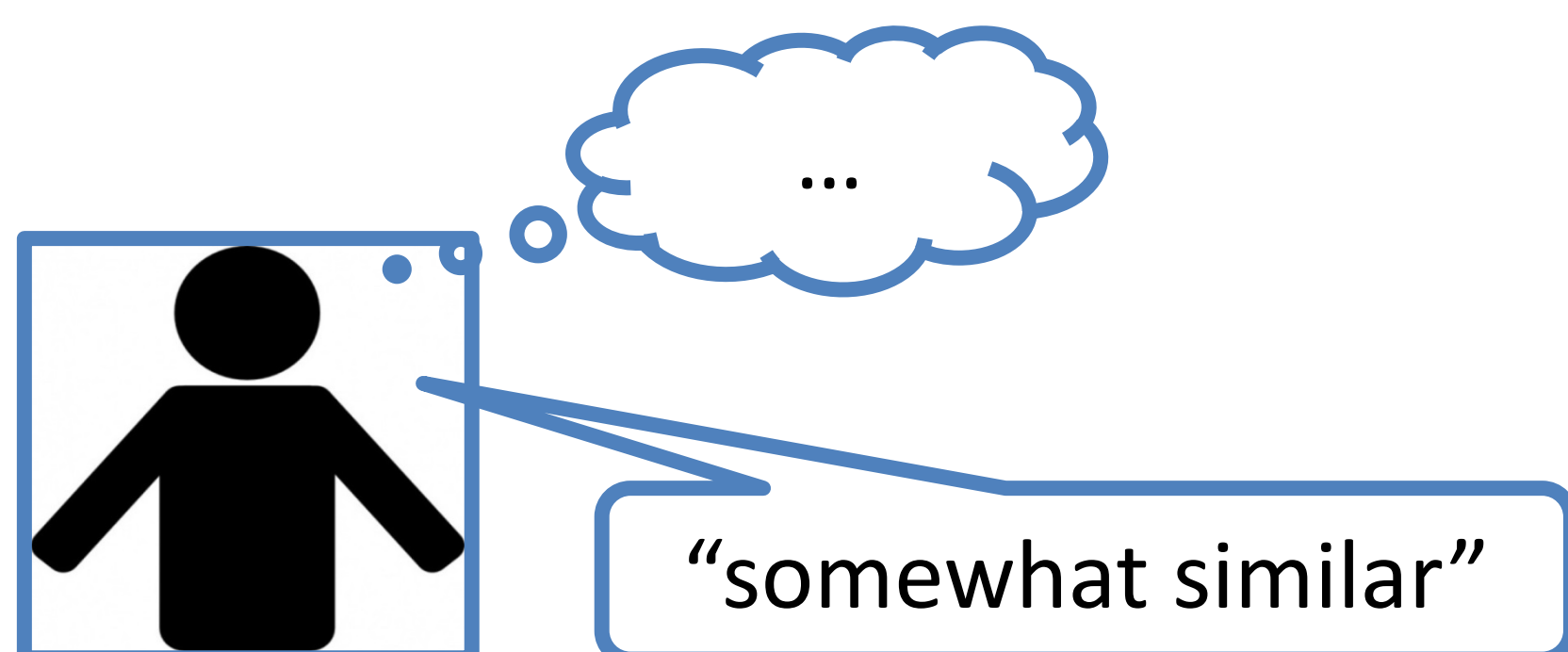
Background. Recent advances in the growing domain of automated driving suggest the need for thoughtful design of human-computer interaction strategies. For example, human drivers can process scene variability on implicit levels, but automated systems require explicit rule-based judgments of similarity and difference. What level of abstraction an automation uses in its visual perception may mean the difference between effective human-automation communication, or “uncanny valley”-like conflicts leading to problems of automation disuse, misuse, or abuse.

Purpose of study. In the present research, different quantifications (**semantic coding** vs. **computer vision features**) of driving scene-to-scene similarity and difference were compared against **intuitive human judgments** as a reference point for future human-automation interactions.

Methods

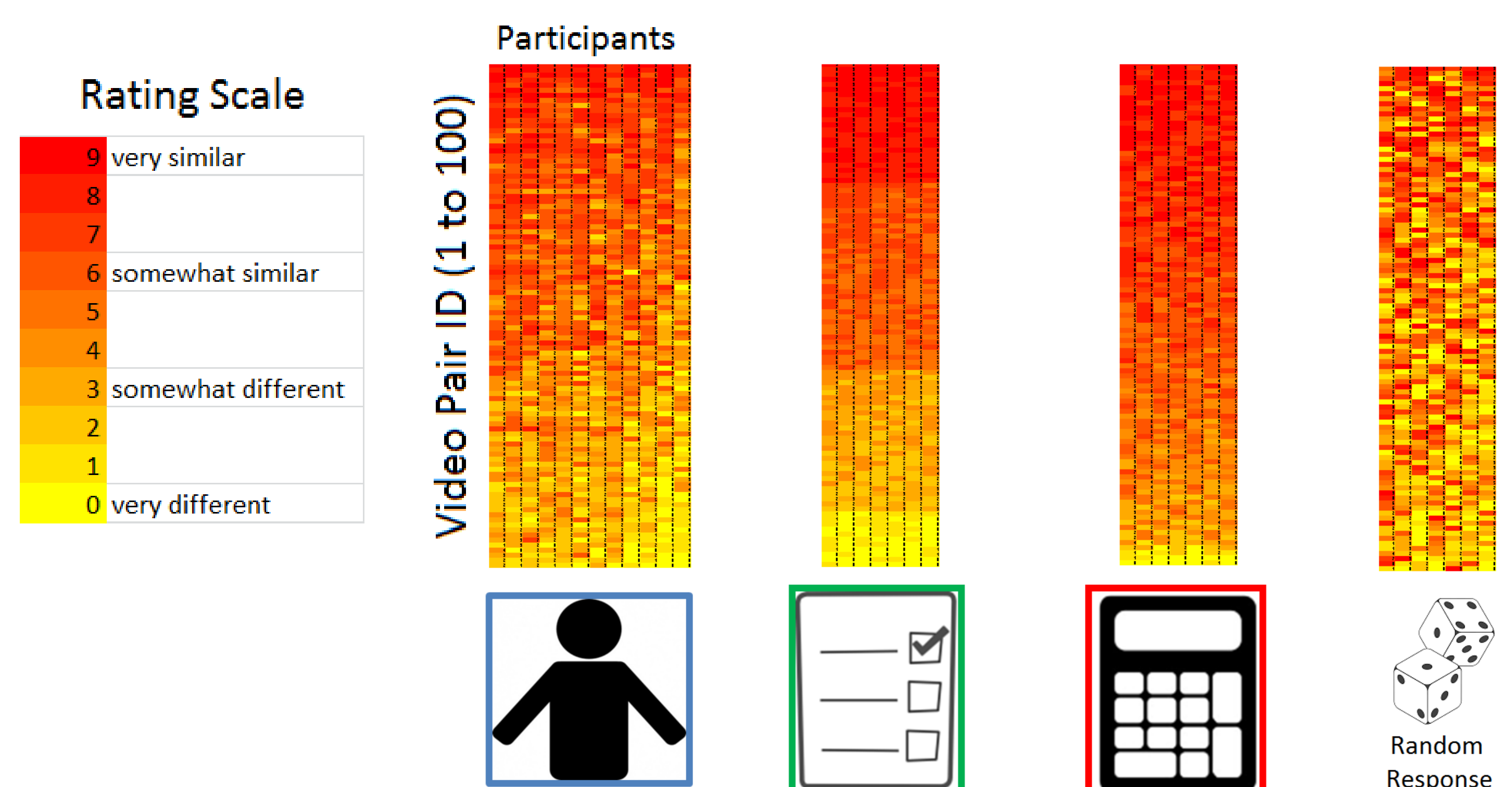
Participants. 12 MSc students (11 male : 1 female)
Mean age = 22.9 yrs old (SD = 1.4)
Mean driving license = 4.8 yrs (SD = 1.9)

Procedure. Each participant rated the same 100 randomly paired driving video clips (i.e., 3 seconds long) on a scale from “0 – Very Different” to “9 – Very Similar”



Results/Conclusions

Scene similarity/difference ratings from **semantic coding** quantification showed **closer matches** to **human participant** judgments than those generated from **computer vision**.



Humans evidence apparent non-random individual differences in judging various driving scenes. Both ‘meaning’ and particularly ‘feature’ level descriptions require improvements to coordinate common ground with human intuition of driving scene similarity/difference.