

LV 706.046 Summer Term 2019
Monday, March, 18

From measuring Usability to measuring Causability: Methods of Explainable AI

Assoc.Prof. Dr. Andreas Holzinger

Holzinger-Group, HCI-KDD, Institute for Medical Informatics/Statistics
Medical University Graz, Austria
&
Institute of interactive Systems & Data Science
Graz University of Technology, Austria

a.holzinger@tugraz.at

<https://hci-kdd.org/intelligent-user-interfaces-2019>



Holzinger Group hci-kdd.org

1

706.046 AK HCI Summer 2019 L01

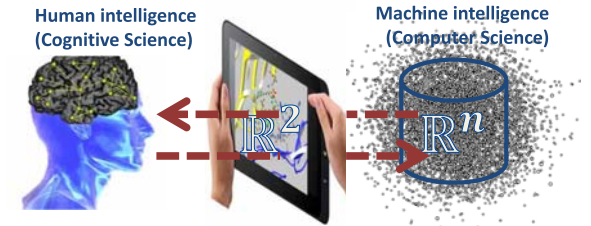
Our Goal in this AK: design, develop & test a System Causability Scale

Holzinger Group hci-kdd.org

2

706.046 AK HCI Summer 2019 L01

- Causability := a property of a person (Human)
- Explainability := a property of a system (Computer)



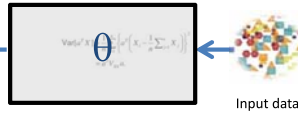
Holzinger Group hci-kdd.org

3

706.046 AK HCI Summer 2019 L01

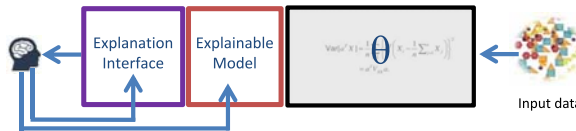
The challenge HCI-KDD

Why did the algorithm do that?
Can I trust these results?
How can I correct an error?



Input data

A possible solution



Input data

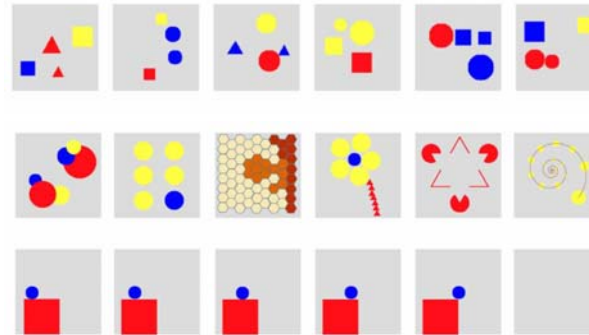
The domain expert can understand why ...
The domain expert can learn and correct errors ...
The domain expert can re-enact on demand ...

Holzinger Group hci-kdd.org

4

706.046 AK HCI Summer 2019 L01

Our Playground HCI-KDD



Holzinger Group hci-kdd.org

5

706.046 AK HCI Summer 2019 L01

First Homework for you: make yourself familiar ... HCI-KDD

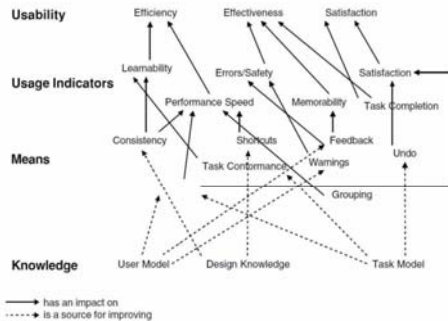
- <https://www.tensorflow.org/tutorials>

Holzinger Group hci-kdd.org

6

706.046 AK HCI Summer 2019 L01

Usability HCI-KDD



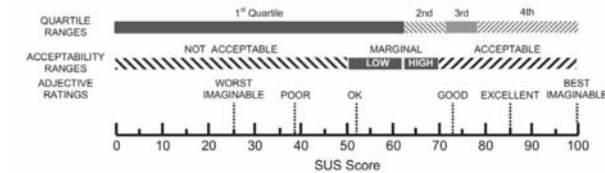
Veer, G. C. v. d. & Welie, M. v. (2004) DUTCH: Designing for Users and Tasks from Concepts to Handles. In: Diaper, D. & Stanton, N. (Eds.) *The Handbook of Task Analysis for Human-Computer Interaction*. Mahwah (New Jersey), Lawrence Erlbaum, 155-173.

Holzinger Group hci-kdd.org

7

706.046 AK HCI Summer 2019 L01

System Usability Scale HCI-KDD



Bangor, A., Kortum, P. T. & Miller, J. T. (2008) An empirical evaluation of the System Usability Scale. *International Journal of Human-Computer Interaction*, 24, 6, 574-594.

Holzinger Group hci-kdd.org

8

706.046 AK HCI Summer 2019 L01

The System Usability Scale Standard Version HCI-KDD

The System Usability Scale Standard Version		Strongly Disagree					Strongly Agree				
		1	2	3	4	5					
1	I think that I would like to use this system frequently.		0	0	0	0	0				
2	I found the system unnecessarily complex.		0	0	0	0	0				
3	I thought the system was easy to use.		0	0	0	0	0				
4	I think that I would need the support of a technical person to be able to use this system.		0	0	0	0	0				
5	I found the various functions in this system were well integrated.		0	0	0	0	0				
6	I thought there was too much inconsistency in this system.		0	0	0	0	0				
7	I would imagine that most people would learn to use this system very quickly.		0	0	0	0	0				
8	I found the system very awkward to use.		0	0	0	0	0				
9	I felt very confident using the system.		0	0	0	0	0				
10	I needed to learn a lot of things before I could get going with this system.		0	0	0	0	0				

Holzinger Group hci-kdd.org

9

706.046 AK HCI Summer 2019 L01

Explainability

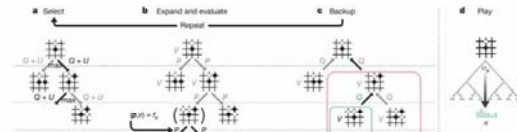


Figure 2 | MCTS in AlphaGo. a, Each simulation traverses the tree by selecting the edge with maximum action value Q plus an upper confidence bound U that depends on a stored prior probability P and visit count N for that edge (which is incremented once traversed). b, The leaf node is expanded and the associated position is evaluated by the neural network $P(x, y, U) = f(x, y)$; the vector of P values are stored in the outgoing edges from x . c, Action value Q is updated to track the mean of all evaluations V in the subtree below that action. & Once the search is complete, search probabilities π are returned, proportional to N^{π} , where N is the visit count of each move from the root state and v is a parameter controlling temperature.

19 OCTOBER 2017 | VOL. 550 | NATURE | 355

$$(p, v) = f_{\theta}(s) \text{ and } l = (z - v)^2 - \pi^T \log p + c \|\theta\|^2$$

David Silver, Julian Schrittwieser, Karen Simonyan, Ioannis Antonoglou, Aja Huang, Arthur Guez, Thomas Hubert, Lucas Baker, Matthew Lai, Adrian Bolton, Yutian Chen, Timothy Lillicrap, Fan Hui, Laurent Sifre, George Van Den Driessche, Thore Graepel & Demis Hassabis 2017. Mastering the game of Go without human knowledge. Nature, 550, (7676), 354-359, doi:10.1038/nature24270.



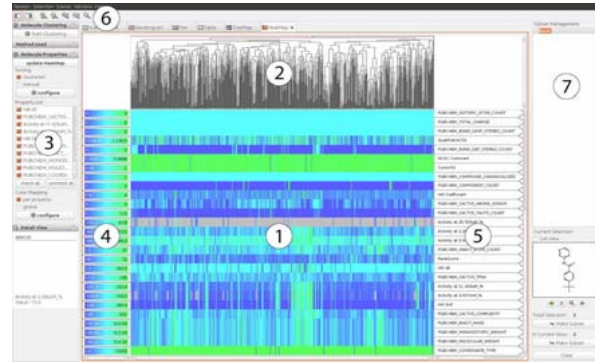
David Silver, Aja Huang, Chris J. Maddison, Arthur Guez, Laurent Sifre, George Van Den Driessche, Julian Schrittwieser, Ioannis Antonoglou, Veda Panneershelvam, Marc Lanctot, Sander Dieleman, Dominik Grewe, John Nham, Nal Kalchbrenner, Ilya Sutskever, Timothy Lillicrap, Madeleine Leach, Koray Kavukcuoglu, Thore Graepel & Demis Hassabis 2016. Mastering the game of Go with deep neural networks and tree search. Nature, 529, (7587), 484-489, doi:10.1038/nature16961.

Example for an Explanation Interface



Todd Kulesza, Margaret Burnett, Weng-Keen Wong & Simone Stumpf. Principles of explanatory debugging to personalize interactive machine learning. Proceedings of the 20th International Conference on Intelligent User Interfaces (UII 2015), 2015 Atlanta. ACM, 126-137, doi:10.1145/2678025.2701399.

Example for an Explanation Interface - open work

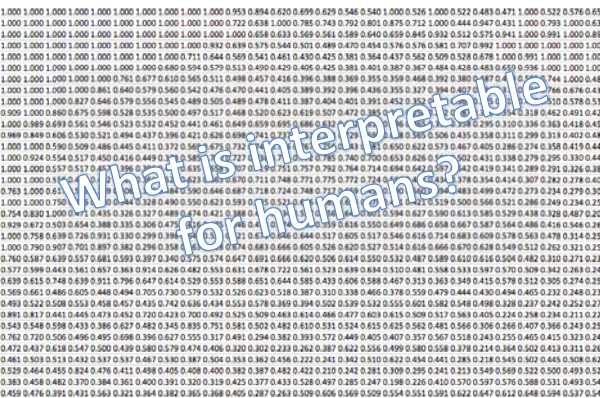


Werner Sturm, Till Schaefer, Tobias Schreck, Andreas Holzinger & Torsten Ullrich. Extending the Scaffold Hunter Visualization Toolkit with Interactive Heatmaps In: Borgo, Rita & Turkey, Cagatay, eds. EG UK Computer Graphics & Visual Computing CGVC 2015, 2015 University College London (UCL). Euro Graphics (EG), 77-84, doi:10.2312/cgvc.20151247.

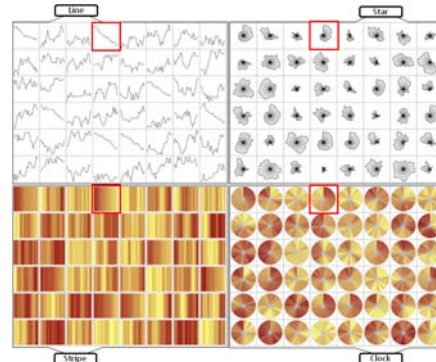
Research and Development in Interpretability ...

- Interpretability as a novel kind for supporting teaching, learning and knowledge discovery,
- Particularly in abstract fields (informatics)
- Compliance to European Law “the right of explanation”
- Check for bias in machine learning results
- Fostering trust, acceptance, making clear the reliability

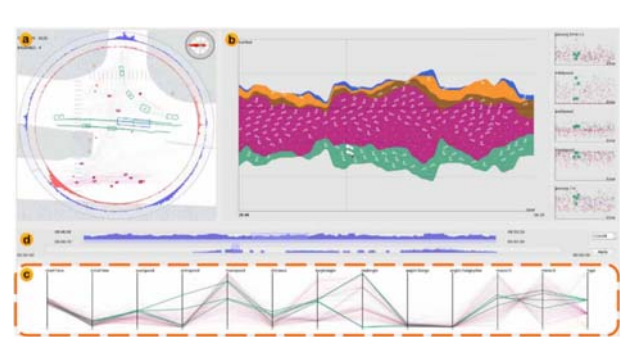
Andreas Holzinger 2018. Explainable AI (ex-AI). Informatik-Spektrum, 41, (2), 138-143, doi:10.1007/s00287-018-1102-5.



What is understandable, interpretable, intelligible?



Explainable AI is a huge challenge for visualization



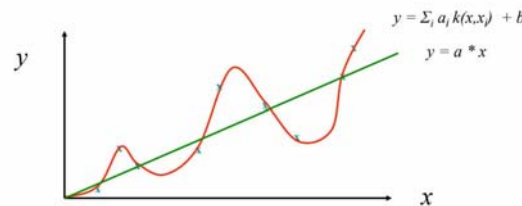
Methods of Explainable AI

- 1) Gradients
- 2) Sensitivity Analysis
- 3) Decomposition Relevance Propagation (Pixel-RP, Layer-RP, Deep Taylor Decomposition, ...)
- 4) Optimization (Local-IME – model agnostic, BETA transparent approximation, ...)
- 5) Deconvolution and Guided Backpropagation
- 6) Model Understanding
 - Feature visualization, Inverting CNN
 - Qualitative Testing with Concept Activation Vectors TCAV
 - Network Dissection

Andreas Holzinger LV 706.315 From explainable AI to Causability, 3 ECTS course at Graz University of Technology
<https://hci-kdd.org/explainable-ai-causability-2019> (course given since 2016)

- **Deductive Reasoning** = Hypothesis > Observations > Logical Conclusions
 - DANGER: Hypothesis must be correct! DR defines whether the truth of a conclusion can be determined for that rule, based on the truth of premises: $A=B$, $B=C$, therefore $A=C$
- **Inductive reasoning** = makes broad generalizations from specific observations
 - DANGER: allows a conclusion to be false if the premises are true
 - generate hypotheses and use DR for answering specific questions
- **Abductive reasoning** = inference = to get the best explanation from an incomplete set of preconditions.
 - Given a true conclusion and a rule, it attempts to select some possible premises that, if true also, may support the conclusion, though not uniquely.
 - Example: "When it rains, the grass gets wet. The grass is wet. Therefore, it might have rained." This kind of reasoning can be used to develop a hypothesis, which in turn can be tested by additional reasoning or data.

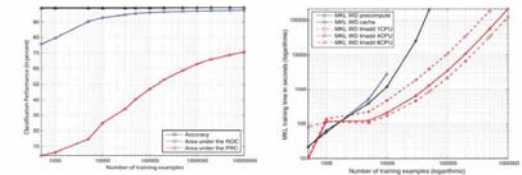
- := information provided by direct observation (empirical evidence) in contrast to information provided by inference
 - Empirical evidence = information acquired by observation or by experimentation in order to verify the truth (fit to reality) or falsify (non-fit to reality).
 - Empirical inference = drawing conclusions from empirical data (observations, measurements)
 - Causal inference = drawing a conclusion about a causal connection based on the conditions of the occurrence of an effect.
 - Causal inference is an example of causal reasoning.



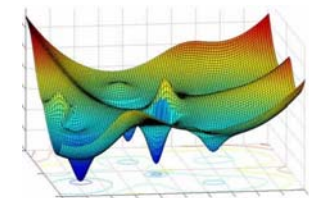
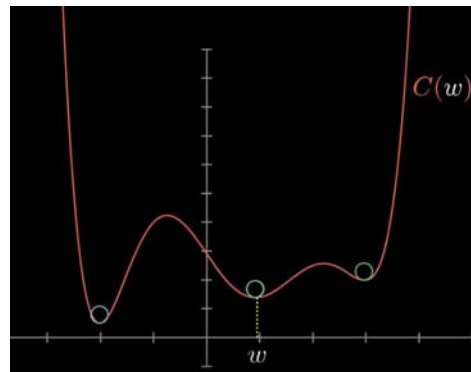
Gottfried W. Leibniz (1646-1716)
 Hermann Weyl (1885-1955)
 Vladimir Vapnik (1936-)
 Alexey Chervonenkis (1938-2014)
 Gregory Chaitin (1947-)

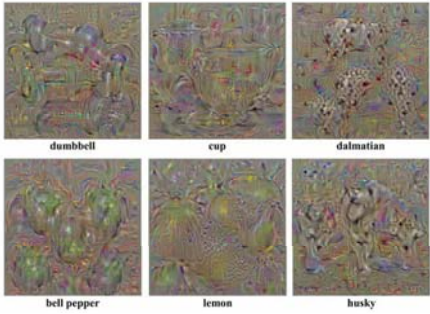


- High dimensionality (curse of dim., many factors contribute)
- Complexity (real-world is non-linear, non-stationary, non-IID *)
- Need of large top-quality data sets
- Little prior data (no mechanistic models of the data)
 - *) = Def.: a sequence or collection of random variables is independent and identically distributed if each random variable has the same probability distribution as the others and all are mutually independent



Sören Sonnenburg, Gunnar Rätsch, Christin Schaefer & Bernhard Schölkopf 2006. Large scale multiple kernel learning. Journal of Machine Learning Research, 7, (7), 1531-1565.





Karen Simonyan, Andrea Vedaldi & Andrew Zisserman 2013. Deep inside convolutional networks: Visualising image classification models and saliency maps. arXiv:1312.6034.

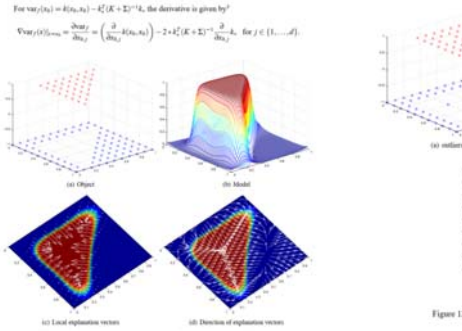


Figure 1: Explaining simple object classification with Gaussian Processes. Panel (a) of Figure 1 shows the training data of a simple object classification task and panel (b) shows the model learned using GP. The data is labeled -1 for the blue points and +1 for the red points. As illustrated in panel (b) the model is a probability function for the positive class which gives every data point a probability of being in this class. Panel (c) shows the probability gradient of the model together with the local Gaussian explanation vectors. Along the hypotenuse and at the corners of the triangle explanations from both features interact towards the triangle class while along

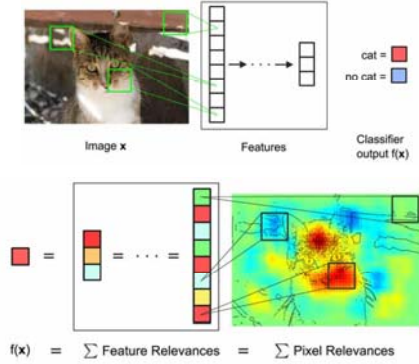
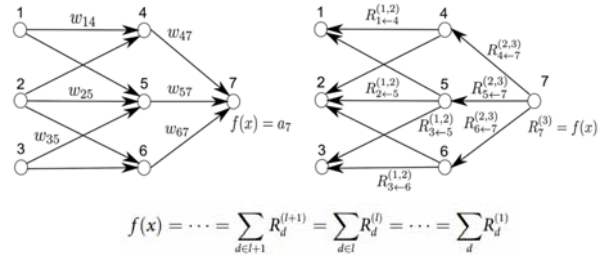


Figure 1: Explaining simple object classification with Gaussian Processes. Panel (a) of Figure 1 shows the training data of a simple object classification task and panel (b) shows the model learned using GP. The data is labeled -1 for the blue points and +1 for the red points. As illustrated in panel (b) the model is a probability function for the positive class which gives every data point a probability of being in this class. Panel (c) shows the probability gradient of the model together with the local Gaussian explanation vectors. Along the hypotenuse and at the corners of the triangle explanations from both features interact towards the triangle class while along



Sebastian Bach, Alexander Binder, Grégoire Montavon, Frederick Klauschen, Klaus-Robert Müller & Wojciech Samek 2015. On pixel-wise explanations for non-linear classifier decisions by layer-wise relevance propagation. PLoS one, 10, (7), e0130140, doi:10.1371/journal.pone.0130140.

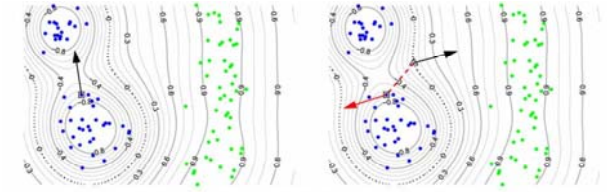


Figure 1: Explaining simple object classification with Gaussian Processes. Panel (a) of Figure 1 shows the training data of a simple object classification task and panel (b) shows the model learned using GP. The data is labeled -1 for the blue points and +1 for the red points. As illustrated in panel (b) the model is a probability function for the positive class which gives every data point a probability of being in this class. Panel (c) shows the probability gradient of the model together with the local Gaussian explanation vectors. Along the hypotenuse and at the corners of the triangle explanations from both features interact towards the triangle class while along

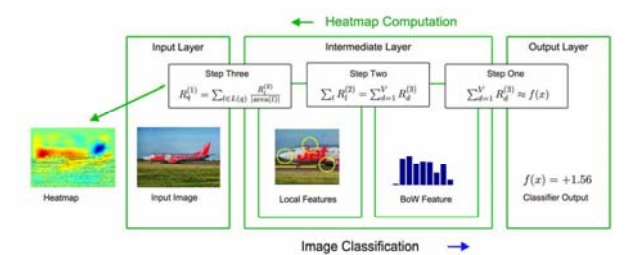
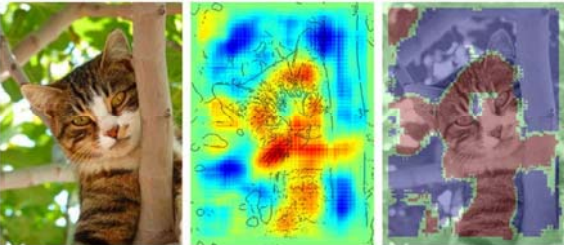


Figure 1: Explaining simple object classification with Gaussian Processes. Panel (a) of Figure 1 shows the training data of a simple object classification task and panel (b) shows the model learned using GP. The data is labeled -1 for the blue points and +1 for the red points. As illustrated in panel (b) the model is a probability function for the positive class which gives every data point a probability of being in this class. Panel (c) shows the probability gradient of the model together with the local Gaussian explanation vectors. Along the hypotenuse and at the corners of the triangle explanations from both features interact towards the triangle class while along



Definition 1. A heatmapping $R(x)$ is *conservative* if the sum of assigned relevances in the pixel space corresponds to the total relevance detected by the model:

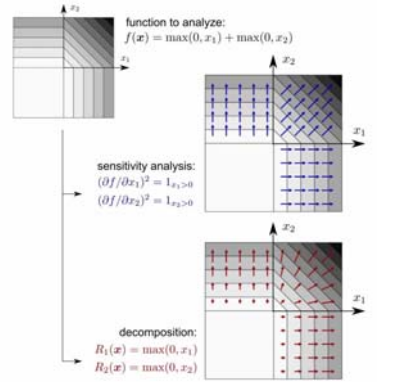
$$\forall x: f(x) = \sum_p R_p(x).$$

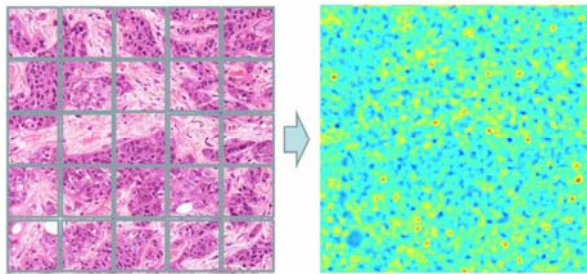
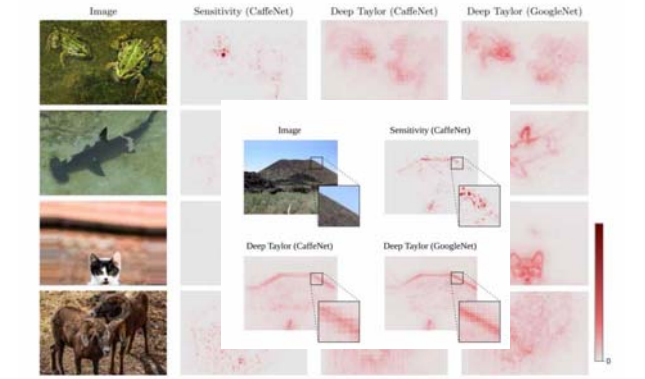
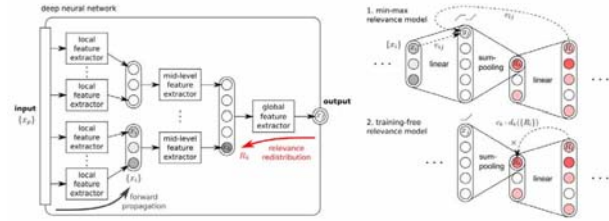
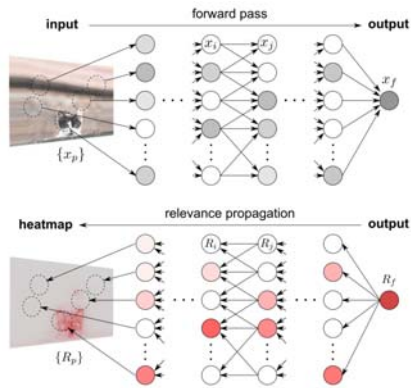
Definition 2. A heatmapping $R(x)$ is *positive* if all values forming the heatmap are greater or equal to zero, that is:

$$\forall x, p: R_p(x) \geq 0$$

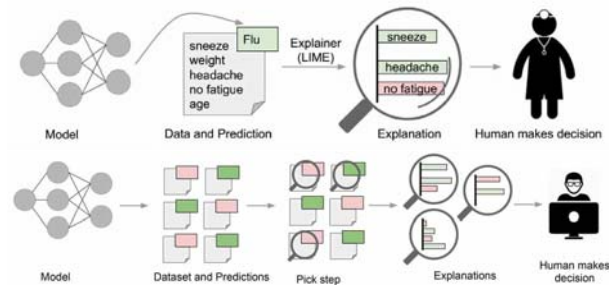
Definition 3. A heatmapping $R(x)$ is *consistent* if it is conservative and positive. That is, it is consistent if it complies with Definitions 1 and 2.

Gregoire Montavon, Sebastian Lapuschkin, Alexander Binder, Wojciech Samek & Klaus-Robert Müller 2017. Explaining nonlinear classification decisions with deep Taylor decomposition. Pattern Recognition, 65, 211-222, doi:10.1016/j.patcog.2016.11.008.

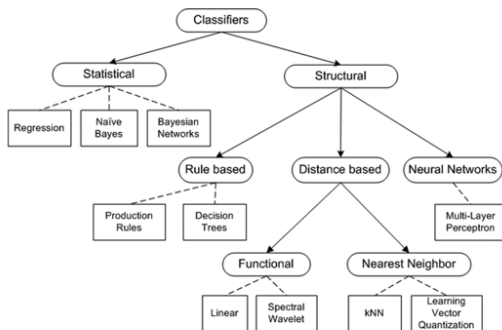




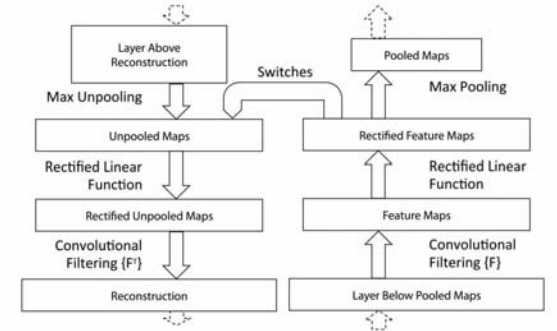
Sebastian Bach, Alexander Binder, Grégoire Montavon, Frederick Klauschen, Klaus-Robert Müller & Wojciech Samek 2015. On pixel-wise explanations for non-linear classifier decisions by layer-wise relevance propagation. PLoS one, 10, (7), e0130140, doi:10.1371/journal.pone.0130140.



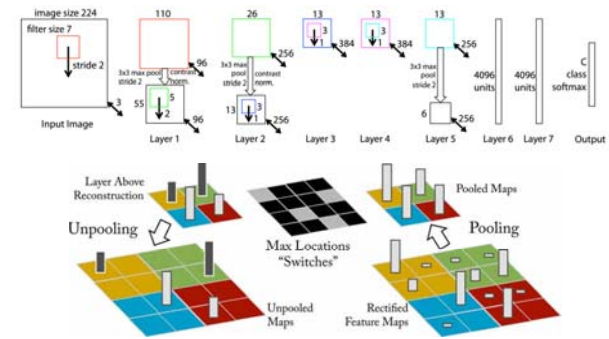
Marco Tulio Ribeiro, Sameer Singh & Carlos Guestrin. Why should I trust you?: Explaining the predictions of any classifier. 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 2016. ACM, 1135-1144, doi:10.1145/2939672.2939778.



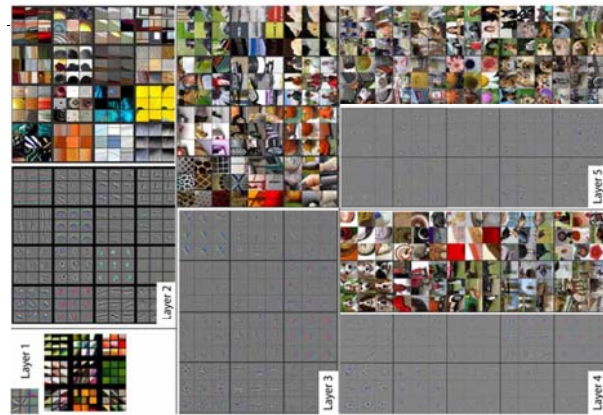
Himabindu Lakkaraju, Ece Kamar, Rich Caruana & Jure Leskovec 2017. Interpretable and Explorable Approximations of Black Box Models. arXiv:1707.01154.



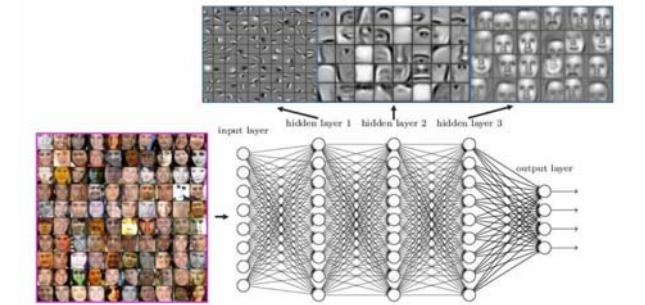
Matthew D. Zeiler & Rob Fergus 2013. Visualizing and Understanding Convolutional Networks. arXiv:1311.2901.



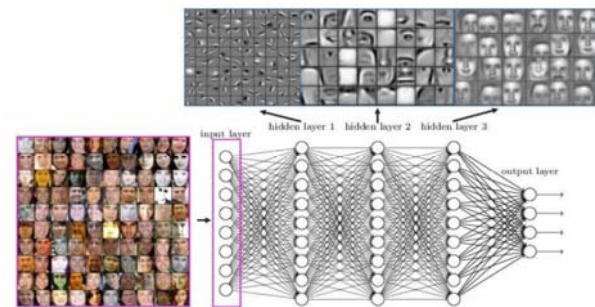
Matthew D. Zeiler & Rob Fergus 2014. Visualizing and understanding convolutional networks. In: D., Fleet, T., Pajdla, B., Schiele & T., Tuytelaars (eds.) ECCV, Lecture Notes in Computer Science LNCS 8689. Cham: Springer, pp. 818-833, doi:10.1007/978-3-319-10590-1_53.
Holzinger Group hci-kdd.org 46 706.046 AK HCI Summer 2019 L01



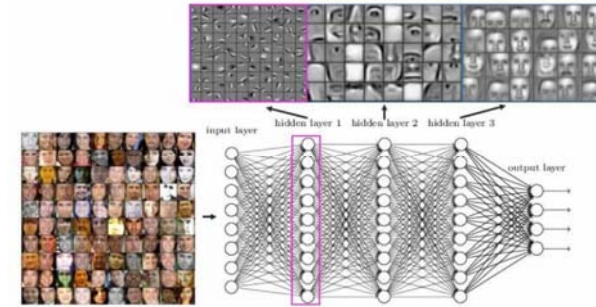
Matthew D. Zeiler & Rob Fergus 2013. Visualizing and Understanding Convolutional Networks. arXiv:1311.2901.
Holzinger Group hci-kdd.org 47 706.046 AK HCI Summer 2019 L01



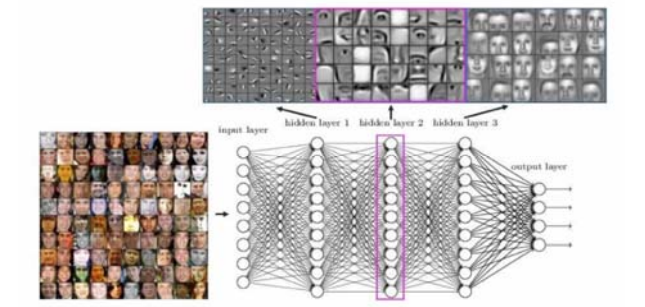
Matthew D. Zeiler & Rob Fergus 2013. Visualizing and Understanding Convolutional Networks. arXiv:1311.2901
Holzinger Group hci-kdd.org 48 706.046 AK HCI Summer 2019 L01



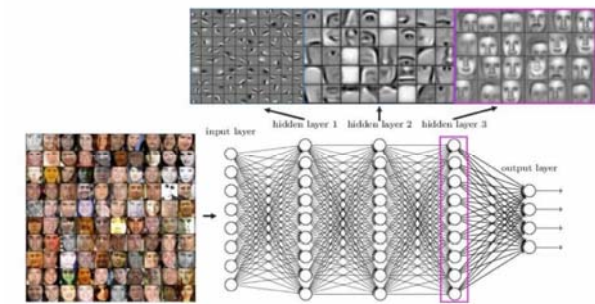
Holzinger Group hci-kdd.org 49 706.046 AK HCI Summer 2019 L01



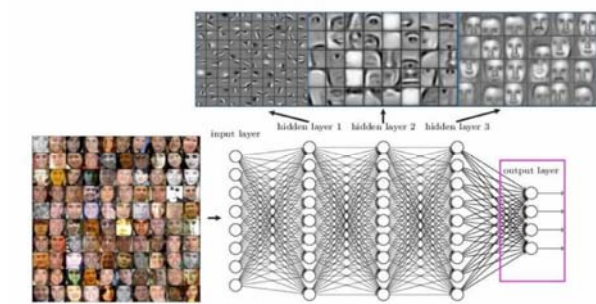
Holzinger Group hci-kdd.org 50 706.046 AK HCI Summer 2019 L01



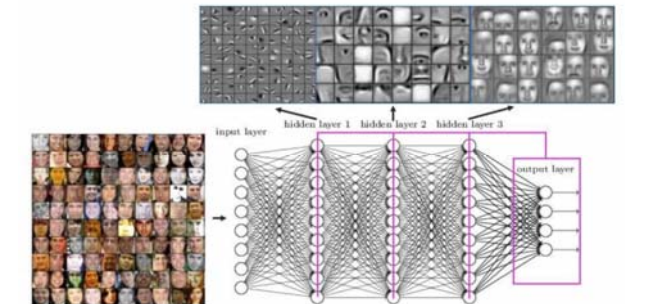
Holzinger Group hci-kdd.org 51 706.046 AK HCI Summer 2019 L01



Holzinger Group hci-kdd.org 52 706.046 AK HCI Summer 2019 L01



Holzinger Group hci-kdd.org 53 706.046 AK HCI Summer 2019 L01



Holzinger Group hci-kdd.org 54 706.046 AK HCI Summer 2019 L01

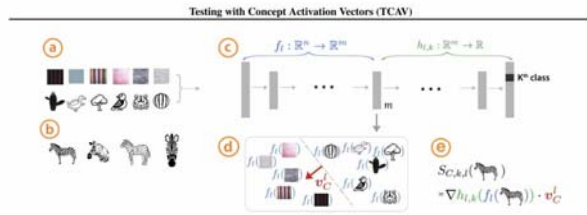


Figure 1. Testing with Concept Activation Vectors: Given a user-defined set of examples for a concept (e.g., "striped"), and random examples (a), labeled training-data examples for the studied class (zebras) (b), and a trained network (c), TCAV can quantify the model's sensitivity to the concept for that class. CAVs are learned by training a linear classifier to distinguish between the activations produced by a concept's examples and examples in any layer (d). The CAV is the vector orthogonal to the classification boundary (v_C , red arrow). For the class of interest (zebras), TCAV uses the directional derivative $SC_{A,k}(x)$ to quantify conceptual sensitivity (e).

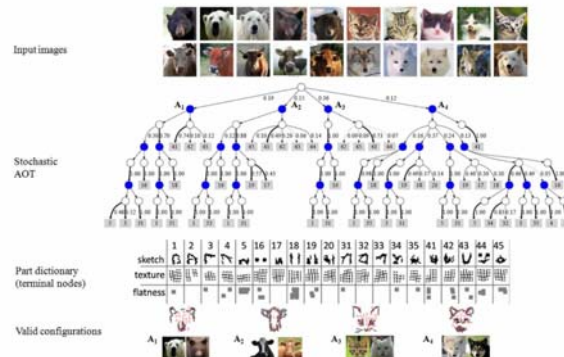
<https://github.com/tensorflow/tcav>

Been Kim, Martin Wattenberg, Justin Gilmer, Carrie Cai, James Wexler & Fernanda Viegas. Interpretability beyond feature attribution: Quantitative testing with concept activation vectors (TCAV). International Conference on Machine Learning (ICML), 2018. 2673-2682.

Holzinger Group hci-kdd.org

55

706.046 AK HCI Summer 2019 L01

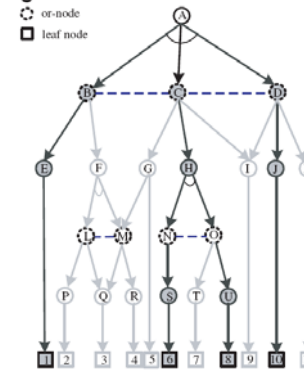


Zhangzhang Si & Song-Chun Zhu 2013. Learning and-or templates for object recognition and detection. IEEE transactions on pattern analysis and machine intelligence, 35, (9), 2189-2205, doi:10.1109/TPAMI.2013.35.

Holzinger Group hci-kdd.org

56

706.046 AK HCI Summer 2019 L01



- Algorithm for this framework
- Top-down/bottom-up computation
- Generalization of small sample
- Use Monte Carlo simulation to synthesis more configurations
- Fill semantic gap

Images credit to Zhaoyin Jia (2009)

Holzinger Group hci-kdd.org

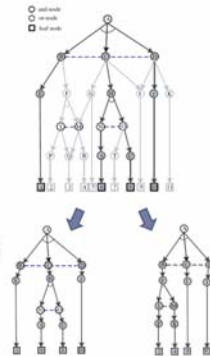
57

706.046 AK HCI Summer 2019 L01

- Terminal (leaf) node: $T(pg)$
- And-Or node: $V^{or}(pg), V^{and}(pg)$
- Set of links: $E(pg)$
- Switch variable at Or-node: $w(t)$
- Attributes of primitives: $\alpha(t)$

$$p(pg; \Theta, R, \Delta) = \frac{1}{Z(\Theta)} \exp(-\xi(pg))$$

$$\xi(pg) = \sum_{v \in V^{or}(pg)} \lambda_v(w(v)) + \sum_{v \in V^{and}(pg) \sim T(pg)} \lambda_t(\alpha(t)) + \sum_{(i,j) \in E(pg)} \lambda_{ij}(v_i, v_j, \gamma_{ij}, \rho_{ij})$$



Holzinger Group hci-kdd.org

58

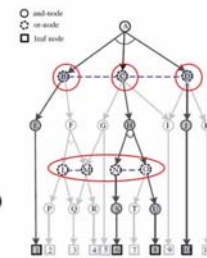
706.046 AK HCI Summer 2019 L01

- Terminal (leaf) node: $T(pg)$
- And-Or node: $V^{or}(pg), V^{and}(pg)$
- Set of links: $E(pg)$
- Switch variable at Or-node: $w(t)$
- Attributes of primitives: $\alpha(t)$

$$p(pg; \Theta, R, \Delta) = \frac{1}{Z(\Theta)} \exp(-\xi(pg))$$

$$\xi(pg) = \sum_{v \in V^{or}(pg)} \lambda_v(w(v)) + \sum_{v \in V^{and}(pg) \sim T(pg)} \lambda_t(\alpha(t)) + \sum_{(i,j) \in E(pg)} \lambda_{ij}(v_i, v_j, \gamma_{ij}, \rho_{ij})$$

SCFG: weigh the frequency at the children of or-nodes



Holzinger Group hci-kdd.org

59

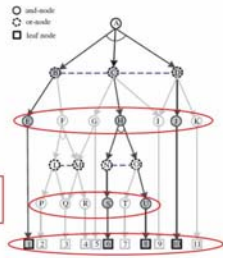
706.046 AK HCI Summer 2019 L01

- Terminal (leaf) node: $T(pg)$
- And-Or node: $V^{or}(pg), V^{and}(pg)$
- Set of links: $E(pg)$
- Switch variable at Or-node: $w(t)$
- Attributes of primitives: $\alpha(t)$

$$p(pg; \Theta, R, \Delta) = \frac{1}{Z(\Theta)} \exp(-\xi(pg))$$

$$\xi(pg) = \sum_{v \in V^{or}(pg)} \lambda_v(w(v)) + \sum_{v \in V^{and}(pg) \sim T(pg)} \lambda_t(\alpha(t)) + \sum_{(i,j) \in E(pg)} \lambda_{ij}(v_i, v_j, \gamma_{ij}, \rho_{ij})$$

Weigh the local compatibility of primitives (geometric and appearance)



Holzinger Group hci-kdd.org

60

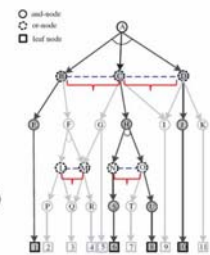
706.046 AK HCI Summer 2019 L01

- Terminal (leaf) node: $T(pg)$
- And-Or node: $V^{or}(pg), V^{and}(pg)$
- Set of links: $E(pg)$
- Switch variable at Or-node: $w(t)$
- Attributes of primitives: $\alpha(t)$

$$p(pg; \Theta, R, \Delta) = \frac{1}{Z(\Theta)} \exp(-\xi(pg))$$

$$\xi(pg) = \sum_{v \in V^{or}(pg)} \lambda_v(w(v)) + \sum_{v \in V^{and}(pg) \sim T(pg)} \lambda_t(\alpha(t)) + \sum_{(i,j) \in E(pg)} \lambda_{ij}(v_i, v_j, \gamma_{ij}, \rho_{ij})$$

Spatial and appearance between primitives (parts or objects)



Holzinger Group hci-kdd.org

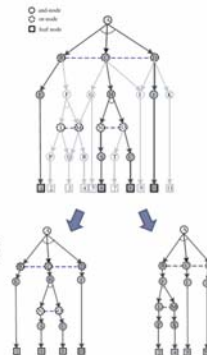
61

706.046 AK HCI Summer 2019 L01

- Terminal (leaf) node: $T(pg)$
- And-Or node: $V^{or}(pg), V^{and}(pg)$
- Set of links: $E(pg)$
- Switch variable at Or-node: $w(t)$
- Attributes of primitives: $\alpha(t)$

$$p(pg; \Theta, R, \Delta) = \frac{1}{Z(\Theta)} \exp(-\xi(pg))$$

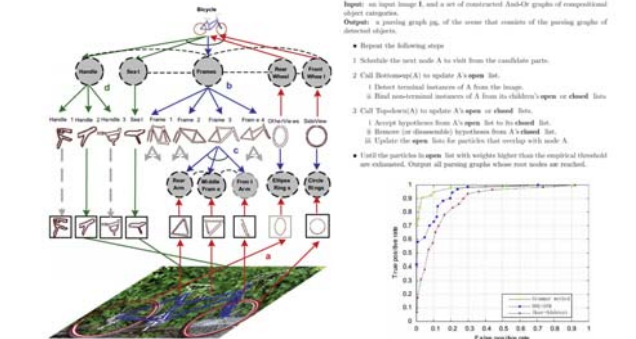
$$\xi(pg) = \sum_{v \in V^{or}(pg)} \lambda_v(w(v)) + \sum_{v \in V^{and}(pg) \sim T(pg)} \lambda_t(\alpha(t)) + \sum_{(i,j) \in E(pg)} \lambda_{ij}(v_i, v_j, \gamma_{ij}, \rho_{ij})$$



Holzinger Group hci-kdd.org

62

706.046 AK HCI Summer 2019 L01



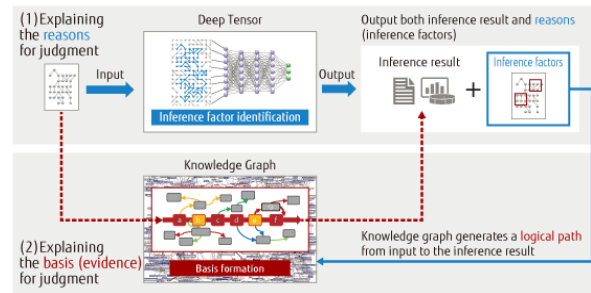
Liang Lin, Tianfu Wu, Jake Porway & Zijian Xu 2009. A stochastic graph grammar for compositional object representation and recognition. Pattern Recognition, 42, (7), 1297-1307, doi:10.1016/j.patcog.2008.10.033.

Holzinger Group hci-kdd.org

63

706.046 AK HCI Summer 2019 L01

Future Work



Explainable AI with Deep Tensor and Knowledge Graph

http://www.fujitsu.com/jp/Images/artificial-intelligence-en_tcm102-3781779.png

- What is a good explanation?
- (obviously if the other did understand it)
- Experiments needed!
- What is explainable/understandable/intelligible?
- When is it enough (Sättigungsgrad – you don't need more explanations – enough is enough)
- But how much is it ...



<https://www.newyorker.com/cartoon/a19697>

Neel C.F. Codella,* Michael Hind,* Karthikeyan Natesan Ramamurthy, Murray Campbell, Amit Dhurandhar, Kush R. Varshney, Dennis Wei & Aleksandra Mojsilovic 2018. Teaching Meaningful Explanations. arXiv:1805.11648.

iv:1805.11648v1 [cs.AI] 29 May 2018

Teaching Meaningful Explanations

Neel C. F. Codella,* Michael Hind,* Karthikeyan Natesan Ramamurthy,* Murray Campbell, Amit Dhurandhar, Kush R. Varshney, Dennis Wei, Aleksandra Mojsilovic*

* These authors contributed equally.

IBM Research
Yorktown Heights, NY 10598
{nccodell,hindh,natesan,mcamp,adharan,kvvarshn,dwei,alekmoan}@us.ibm.com

Abstract

The adoption of machine learning in high-stakes applications such as healthcare and law has lagged in part because predictions are not accompanied by explanations comprehensible to the domain user, who often holds ultimate responsibility for decisions and outcomes. In this paper, we propose an approach to generate such explanations in which training data is augmented to include, in addition to features and labels, explanations elicited from domain users. A joint model is then learned to produce both labels and explanations from the input features. This simple idea ensures that explanations are tailored to the complexity expectations and domain knowledge of the consumer. Evaluation spans multiple modeling techniques on a simple game dataset, an image dataset, and a chemical order dataset, showing that our approach is generalizable across domains and algorithms. Results demonstrate that meaningful explanations can be reliably taught to machine learning algorithms, and in some cases, improve modeling accuracy.

1 Introduction

New regulations call for automated decision making systems to provide “meaningful information” on the logic used to reach conclusions [1–4]. Selbst and Powles interpret the concept of “meaningful information” as information that should be understandable to the audience (potentially individuals