

Tabla de Contenidos

Causalidad y redes neuronales 1

Causalidad y redes neuronales

The extreme case of model interpretability is when we are trying to establish a mechanistic model, that is, a model that actually captures the phenomena behind the data. Good examples include trying to guess whether two molecules (e.g. drugs, proteins, nucleic acids, etc.) interact in a particular cellular environment or hypothesizing how a particular marketing strategy is having an actual effect on sales. Nothing really beats old-style Bayesian methods informed by expert opinion in this realm; they are our best (if imperfect) way we have to represent and infer causality. Vicarious has some [nice recent work](#) illustrating why this more principled approach generalizes better than deep learning in videogame tasks.

— <http://hyperparameter.space/blog/when-not-to-use-deep-learning/>

[Deep Convolutional Neural Networks for Pairwise Causality](#)

[What's the relation between hierarchical models, neural networks, graphical models, bayesian networks?](#)

From:

<http://filosofias.es/wiki/> - **filosofias.es**

Permanent link:

http://filosofias.es/wiki/doku.php/proyectos/tfg/casos/redes_neuronales/start?rev=1510103544

Last update: **2017/11/08 01:12**

