

Vortragsfolien zum zweiten Workshop im Anwenderkreis "Transparenz in neuronalen Netzen"

Vielen Dank an die Vortragenden und allen Teilnehmenden für diesen gelungenen Workshop.
Wenn Sie keine Veranstaltung des Anwenderkreises verpassen möchten, dann können Sie sich auf dieser Seite zu unserer Mailing Liste anmelden:

[Transparenz in neuronalen Netzen - Anwenderkreis](#)



ZERTIFIZIERTE KI

Qualität sichern. Fortschritt gestalten.

Regellernen zur Schwachstellenanalyse und Erklärbarkeit von Black-Box Modellen

von Dr. Daniel Becker, Fraunhofer IAIS



FraunhoferIAIS_R...etierbarkeit.pdf

Advantages, properties and limitations of model agnostic explainability techniques in Machine Learning

von Dr. Antoine Gautier, QuantPi



QuantPi_Vortrag.pdf