

A Prototype Implementation Supporting Intelligent Big Data Analysis

Sebastian Bruchhaus^{*1}, Alexander Kreibich¹, Thoralf Reis¹ and Marco X. Bornschlegl¹

Abstract

This paper builds upon an existing formal trustworthiness model by instantiating it in a construction-domain Big Data pipeline and BI prototype for risk forecast in road construction projects. It is technically embedded into a Trustworthy AI based Big Data Management (TAI-BDM) compliant Big Data analysis pipeline and communicated visually to domain experts. Building on the TAI-BDM reference model and an existing conceptual trustworthiness model, a three dimensional trust state space (validity, capability, reproducibility) is instantiated with trust feature functions. These define phase specific trust dimension functions and an overall trustworthiness norm, and integrate these into a business intelligence prototype for road construction planning that realizes these concepts as structured scoring heuristics for the three trust dimensions and a single aggregated trust score. The prototype demonstrates how problem domain specific trust feature functions can be modelled and implemented. Additionally, it shows how trust signals can be surfaced along the AI2VIS4BigData pipeline to make the evolution of trust more intelligible to end users. A qualitative cognitive walkthrough with a domain expert suggests that the proposed modeling and visualization approach improves the comprehensibility of trust development along the AI supported Big Data analytic workflow. This work takes a previously introduced abstract mathematical model of trustworthiness in AI based Big Data analytics and demonstrates how the approach can be transferred to other data intensive and safety-critical scientific fields, such as biomedicine.

The main contribution is a domain-specific instantiation of the TAI-BDM trust space, including concrete trust features, phase-specific aggregation functions, and a working proof of concept implementation that exposes trust trajectories to end users.

Keywords

Trustworthy artificial intelligence, Big data analytics, Visual analytics, Business intelligence systems, Trust modeling and measurement, Construction risk forecasting

1. Introduction and Motivation of Trust in Intelligent Big Data Analytics

Artificial Intelligence (AI) can help empower users in visual Big Data analytics, as captured by the AI2VIS4BigData model [1]. Establishing systematic notions and measurements of trustworthiness provides objective evidence for when AI systems may safely move from prototypes into mission-critical operations [2]. Bornschlegl's Trustworthy AI-based Big Data Management (TAI-BDM) model introduces the trust dimensions *validity*, *capability*, and *reproducibility*, on which Bruchhaus et al. build a conceptual mathematical model [3] for Trustworthy AI-based Big Data analysis (TAI2VIS4BigData). They define a three-dimensional trust state space with mapping functions for the effects of each data manipulation and analysis step, complemented by a scalar trustworthiness norm yielding an overall trust score. While this precise mathematical model of trust exists, there is currently no domain-specific instantiation with concrete features, weights, and a running system that demonstrates its use in construction risk BI.

This work addresses three interrelated research problems. First, construction-specific trust features must be defined and quantified to provide evidence for the validity, capability, and reproducibility of each AI2VIS4BigData phase. Second, the abstract mappings that model trust dynamics require instantiation in mathematically grounded feature functions consistent with domain expertise. Finally, an integrated implementation is needed to expose trust states and trajectories through a visual BI interface that enables users to interpret trust information in production-ready AI solutions.

TRIIVIS'26: Workshop on Trust and Reproducibility in Interactive Visual Analytics, June 03–06, 2026, Bari, Italy

^{*}Corresponding author.

✉ sebastian.bruchhaus@fernuni-hagen.de (S. Bruchhaus^{*}); alexander.kreibich@fernuni-hagen.de (A. Kreibich); thoralf.reis@fernuni-hagen.de (T. Reis); marco.bornschlegl@fernuni-hagen.de (M. X. Bornschlegl)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

From these challenges emerge three research questions:

(RQ1) How can TAI-BDM trust dimensions be represented in a machine-readable form for Big Data in construction?

(RQ2) How can data manipulation and analysis steps along the AI2VIS4BigData pipeline affect trust over time?

(RQ3) How can an overall measure of trust be derived and visualized for end users of a Big Data analytics application.

To answer these questions, the TAI-BDM trust space is instantiated with construction-specific trust features for each pipeline phase. Phase-specific feature functions then aggregate these features into normalized values for the three dimensions of trust. The resulting constructs are embedded in a TAI-BDM-compliant prototype with a configurable trust bus and a domain-specific data simulator that generates project data with controllable trust characteristics. A qualitative expert walkthrough assesses whether the resulting trust trajectories and visualizations are intelligible, align with expert intuitions, and suggest refinements for the mathematical model and its parameters.

2. Related State of the Art

In the following, IVIS4BigData provides the process and interaction backbone on which TAI-BDM builds its trust modeling, modeling Big Data analysis as a user-empowering, interactive pipeline from data integration to the consumption of visual views with clearly defined transformation stages, user roles, and configuration activities.

Research on trustworthy AI in data-intensive environments spans architectural reference models, trust-centered data management, and formalizations of trustworthiness in AI pipelines. IVIS4BigData and its successor **AI2VIS4BigData** [4] conceptualize AI-supported visual Big Data analysis as an interactive, user-empowering pipeline from data integration to visual effectuation, with clearly defined phases, kernel algorithms, and artifacts that structure how users steer and interpret analytics workflows. This work adopts AI2VIS4BigData as the process backbone: its ordered phases provide the discrete time steps at which trust-relevant evidence is collected and aggregated into evolving trust states.

On top of this process view, the **TAI-BDM** [2] model introduces a trust-integrated big data management framework that augments AI2VIS4BigData with a dedicated trust bus (c.f. Fig. 1). The trust bus continuously collects metadata such as data quality, model performance and robustness, and explanation signals, and aggregates them along three explicit dimensions—capability, validity, and reproducibility—into an overall notion of system trustworthiness that can be queried and visualized as a “trust trajectory” across the pipeline. Prior work has proposed a conceptual mathematical model for this setting that represents the trust state as a vector in a three-dimensional trust space, with mappings describing how data manipulation and analysis steps update this state and with scalar norms providing an overall trust score [3]. However, existing work remains largely abstract and has not yet demonstrated a domain-specific instantiation with concrete trust features, aggregation functions, and a demonstrator for running BI system for construction risk forecasting, which is the objective of this paper.

Mathematical treatment of trustworthiness in this work builds on a previously introduced conceptual model for TAI-BDM-based data exploration processes, which represents the system’s trust state at any time by a vector $s = (c, v, r)$ in a three-dimensional space $T_M \subset [0, 1]^3$, with coordinates capturing capability, validity, and reproducibility under the current data, model, and environment conditions [3]. An environment variable $U \in [0, 1]$ encodes global conditions (e.g., available expert knowledge) and determines an initial trust state via a mapping $\Phi_0(U) = s_0$, while each data manipulation or analysis step f is modeled by a trustworthiness mapping Φ_f that updates the state and induces a trajectory of trust states over the pipeline [3]. A scalar trustworthiness norm $\|\cdot\| : T_M \rightarrow [0, 1]$ then aggregates each vector-valued state into a normalized trust score subject to basic normalization and monotonicity conditions, so that the completely untrustworthy state $(0, 0, 0)$ attains the minimal value 0 and improvements in any coordinate cannot decrease the score, providing a compact, interpretable summary of overall trustworthiness [3]. Examples of such norms include metric-based constructions,

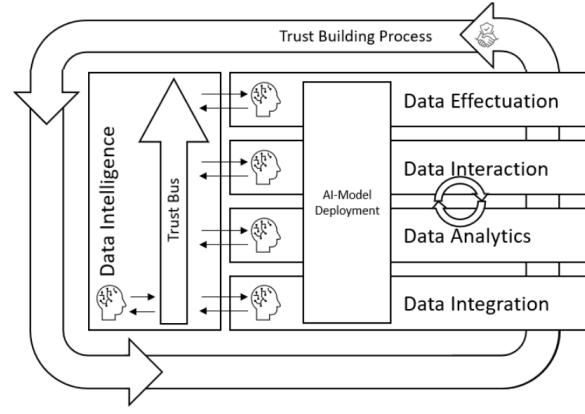


Figure 1: Overview of the TAI-BDM-based big data analysis pipeline with integrated trust bus [2]

weighted sums, and threshold-based variants as discussed in the underlying conceptual model [3]. Relevant candidates are weighted ℓ_p norms, minimum (bottleneck) and maximum norms on $\mathbb{R}^n$, ordered weighted averaging (OWA) operators, and the arithmetic mean as a simple symmetric aggregation function [5].

3. Modeling Trustworthiness in AI-Based Big Data Analysis

The BI prototype addresses a construction-industry use case in which an AI-based analytics pipeline predicts risk and expected profitability of road construction projects in four discrete classes to support early identification of financially critical contracts. In this work, the modeling of trustworthiness is structured in three tightly coupled layers: (1) the AI2VIS4BigData reference model as the analysis and visualization pipeline, (2) the TAI-BDM reference model as a trust-integrated big data management process, and (3) a mathematical formalization of trust states, their dynamics, and an overall trustworthiness measure [1, 2, 3].

3.1. AI2VIS4BigData as Pipeline Context

AI2VIS4BigData provides the conceptual backbone for the analysis workflow, structuring AI-based big data analysis and visualization into a sequence of phases such as data integration, data preparation, data analytics, and visual analysis/effectuation [1]. In order to achieve this, each phase is associated with a set of domain-specific *kernel algorithms* (e.g., for data cleaning, feature engineering, model training, and explanation generation), together with well-defined input and output artifacts for each phase. In the construction-domain use case considered in this paper, these phases are populated with domain-specific operations for road-construction projects, such as the ingestion of project KPIs, generation of derived indicators, training of risk or revenue prediction models, and the visualization of analysis results for decision makers.

Within this layered pipeline, AI2VIS4BigData serves as the “time axis” for the evolution of trust: every execution of a kernel algorithm in a given phase produces not only primary artifacts (e.g., transformed datasets, trained AI models, explanations), but also secondary metadata that can be interpreted as evidence for or against the trustworthiness of the system at that point in the analysis [1, 3]. Therefore, the modeling of trustworthiness takes the AI2VIS4BigData phases and their algorithm sequences as the structural context in which trust features are observed, aggregated, and propagated along the pipeline.

In all four AI2VIS4BigData phases, the same five trust aspects – transparency, correctness, robustness, suitability, and completeness – are measured with suitable indicators at the data, model, and visualization levels. In Data Integration this uses, for example, anomaly scores and distribution comparisons; in Data Analytics, Shapley/LIME values, confusion matrices, and F1/MCC scores. Data Interaction focuses on consistent reporting logic and the suitability of visualizations, while Data Effectuation evaluates

bias between the internal system model and the presented view. All trust feature values are first aggregated via trust dimension functions into reproducibility, validity, and capability, then combined by trustworthiness norms into a single numerical trust score.

3.2. TAI-BDM and the Trust Bus

In the construction BI scenario, domain-specific use contexts (e.g., economic assessment of running road-construction projects, configuration and administration of a BI application) are mapped to concrete TAI-BDM phases [2]. For each AI2VIS4BigData phase inside TAI-BDM, the trust bus is modeled to store partial information on three trust dimensions — reproducibility, validity, and capability — and to update these when kernel algorithms for each phase are executed. This leads to a *trust trajectory* over the pipeline: every step in the TAI-BDM/AI2VIS4BigData process corresponds to a new point in the conceptual trust space T_M , and the trust bus provides the data structures and update rules needed to track this trajectory over time (c.f. Fig. 2). The trust bus thus acts as middleware between the technical AI components and the external stakeholders, enabling, for example, visualization of the current trust state and its history to end users, as well as the use of trust information in governance and control mechanisms [2, 3].

A trust feature function maps technical evidence from an algorithmic step to a numerical trust feature value that reflects how well a specific trust aspect is fulfilled. **Reproducibility features:** These capture whether results can be reliably obtained again under the same conditions. They include explainability, understood as a causal link between input and output, interpretability as the user’s substantive understanding of results, and system transparency as the communication of the current system state. **Validity features:** These focus on the correctness and stability of the processing. Error-freedom denotes correct calibration of models and algorithms, determinism refers to the degree to which identical inputs lead to identical outputs, robustness describes plausible behavior under changed conditions, and impartiality means fair treatment of all relevant subgroups without systematic bias. **Capability features:** These assess how well the system is suited to its task and data. System suitability measures how appropriate the system is for its intended purpose, completeness reflects how fully the system is populated with valid raw or source data, data management concerns the suitability of training data for the applied methods, and performance balance captures the runtime and predictive performance of the algorithms used. **Environmental features:** These describe the initial trustworthiness of the surrounding conditions in which the BI application operates. Cybersecurity covers the protection of the application against cyberattacks, data protection addresses safeguarding users’ digital privacy, AI acceptance reflects the willingness of human staff to cooperate with the AI system, and technology use relates to the avoidance of misuse, such as intrusive IoT tracking or similar practices.

3.3. Mathematical Modeling of Trustworthiness

On top of these architectural models, we recall the mathematical model from [3] and show how we instantiate it via simple, interpretable scoring functions for this domain. This gives a precise semantics to the notion of a “trust state” and its evolution in TAI-BDM-based analyses [3].

Building on this formalism, this work instantiates the abstract trust space and trustworthiness mappings for the specific setting of TAI-BDM-based risk and revenue forecasts in the construction domain [6]. The trust state space T_M is interpreted as the set of all feasible combinations of reproducibility, validity, and capability that can arise from the configured BI application and its data. The environment variable U is specified for each analysis scenario based on domain-expert estimations of global conditions (e.g., data completeness, stability of processes), and the initial state $s_0 = \Phi_0(U)$ is determined accordingly.

For each AI2VIS4BigData phase p , a set of *trust features* is defined that capture phase-specific evidence relevant for the three trust dimensions, such as data-quality metrics in early phases, model performance and robustness indicators in the analytics phase, or visualization and interpretability quality measures in the effectuation phase [6]. Let F_p denote the vector of normalized trust feature values observed in phase p . For every dimension $d \in \{c, v, r\}$ and phase p , a trust dimension function $T_{d,p} : F_p \rightarrow [0, 1]$ is introduced,

which aggregates the feature values into a new value for the corresponding trust dimension [6]. In the present work, these functions are implemented as (possibly weighted) mean operators over feature groups, where the weights are specified and calibrated by domain experts to reflect the relative importance of individual indicators.

Given a current trust state $s = (c, v, r)$ and the feature vector F_p in phase p , the updated state $s' = (c', v', r')$ is obtained by applying the dimension functions $c' = T_{c,p}(F_p)$, $v' = T_{v,p}(F_p)$, $r' = T_{r,p}(F_p)$, so that the phase-specific trustworthiness mapping Φ_{f_p} for the effective manipulation step f_p of phase p is concretely realized via these aggregations and the corresponding parameterization θ_p . In this way, the abstract mappings Φ_f from the conceptual model become implementable functions whose behavior is controlled through configurable feature sets and weights.

Weighted means are chosen because they offer a mathematically transparent, monotone, and easily interpretable aggregation that lets domain experts encode relative importance of heterogeneous trust features while preserving compatibility with existing metric-based norms, which is crucial for safety auditing and fairness discussions.

On top of the vector-valued trust state, the overall trustworthiness norm $\|\cdot\| : T_M \rightarrow [0, 1)$ is instantiated in this case as a variant of the arithmetic mean of the three trust dimensions, optionally combined with a nonlinear transformation to ensure that the resulting values remain strictly within the unit interval and that changes near the boundaries 0 and 1 become progressively harder to achieve. This choice embodies the design decision that capability, validity, and reproducibility should be treated as equally important at the top level, while still allowing differentiated emphasis within the dimension functions $T_{d,p}$ [6].

By repeatedly applying the instantiated mappings Φ_{f_p} along the configured TAI-BDM pipeline, the BI application produces a sequence of trust states $s_0, s_1, \dots, s_n$ and corresponding scalar scores $\|s_0\|, \|s_1\|, \dots, \|s_n\|$, which together form the *trust trajectory* of a given analysis run. This trajectory is stored and exposed by the implemented trust bus and can be visualized in the prototype to provide end users with an interpretable, mathematically well-founded view on how the trustworthiness of the AI-based risk forecasts evolves from raw data to final decisions.

In particular, the trust norm and dimension functions satisfy basic properties of monotone and Lipschitz aggregation operators in ordered metric spaces [5, 7]. For example, for the arithmetic-mean-based trust norm $\|s\| = (c + v + r)/3$, monotonicity follows immediately because if $c' \geq c$, $v' \geq v$, and $r' \geq r$ with at least one strict inequality, then $(c' + v' + r')/3 > (c + v + r)/3$ by strict inequality of real sums.

4. Implementation of an Intelligent Big Data Analytics Workflow

In the implementation phase, the conceptual models are realized as a fully executable TAI-BDM analysis pipeline for construction risk forecasts. The backend uses a Python/Luigi stack [8, 9] that orchestrates configuration, workflow execution, and FAIR-oriented storage of all intermediate and final artefacts for each analysis run [10]. The information model combines a performance assessment subsystem for road-construction KPIs with a trust-bus subsystem, which share a common artefact pool but maintain separate metadata.

The dynamic workflow is implemented as a Luigi-controlled pipeline in which each AI2VIS4Big-Data phase corresponds to a task and a central “trustworthiness file” carries the evolving trust state between phases [9, 8]. The performance subsystem uses GBRT, SVC, and neural-network classifiers from scikit-learn as complementary baselines for tabular risk and revenue prediction and produces SHAP/LIME-based explanation artefacts that are converted into quantitative trust features [11, 12, 13, 14]. For the trust bus, these feature values are computed in each AI2VIS4BigData layer and aggregated by calibrated trust-dimension functions and an arithmetic-mean-based trust norm (with a transformation to keep values in $[0, 1)$) to yield a phase-wise trust trajectory [2, 4].

The User Centered Design (UCD)-oriented frontend is implemented as a Vue.js single-page application with local and Docker deployment variants [15, 16]. It provides forms for configuring and starting

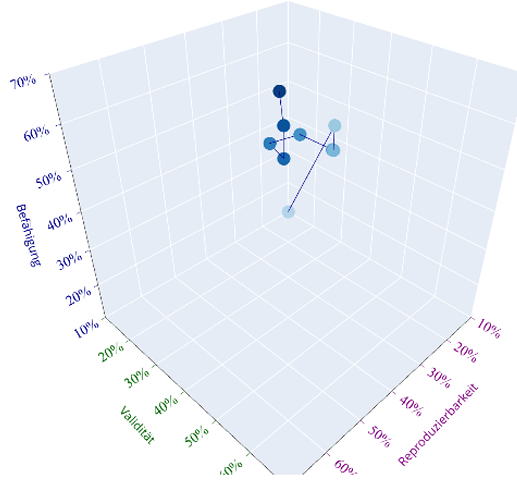


Figure 2: Example trust trajectory along the TAI-BDM pipeline.

new TAI-BDM analyses, views for each AI2VIS4BigData layer, log/console access, and a dedicated visualization of the trust path in the 3D trust space (Fig. 2) [4]. Because no suitable real data were available, a Python-based simulator generates synthetic road-construction projects with controllable data quality, label noise, covariate shift, and sample size, assuming i.i.d. tabular data without temporal or relational structure, which tends to produce smoother, more stable trust trajectories than in real-world construction analytics [8].

5. Evaluation of Trustworthiness

This section evaluates the proposed trust scoring scheme and its implementation in the BI prototype, with a focus on the three-dimensional trust state space, the phase-specific trust dimension functions, and the overall trustworthiness norm. The goal is to assess whether these constructs behave consistently with their intended semantics and whether the resulting trust trajectories are interpretable for expert users.

5.1. Qualitative Walkthrough and Main Observations

Given the exploratory nature of the scoring scheme and the absence of canonical “true” trust trajectories, we conducted a qualitative cognitive walkthrough with a single expert familiar with AI, machine learning, and evaluation concepts. The expert interacted with the running prototype using preconfigured simulator scenarios that varied data quality, model configuration, and pipeline behavior, and was asked to interpret capability, validity, and reproducibility at selected TAI-BDM phases, explain changes in these values in terms of underlying pipeline events (e.g., data cleaning, model training, recalibration), and judge whether the scalar trust scores and their evolution appeared plausible given the scenario descriptions. Across scenarios, the relative behavior of the three trust dimensions largely matched the expert’s expectations: improved data quality and model performance yielded monotonic or non-decreasing capability and validity, while aggressive feature selection, questionable sampling, or opaque models led to plausible decreases or stagnation, especially in validity and reproducibility. The three-dimensional trust view was perceived as helpful to distinguish, for example, states with high capability but lower validity and reproducibility (“works well on this data, but may not generalize”), and the arithmetic-mean trust norm was regarded as coherent with the dimension-wise behavior, with simultaneous improvements in all dimensions producing clear score increases and single-dimension changes having more moderate effects.

5.2. Sensitivity, Weights, and Norm Interpretation

The walkthrough highlighted that the sensitivity of trust scores is strongly governed by the chosen feature weights within each dimension function. Small changes in heavily weighted features could noticeably shift a dimension value, while substantial changes in lightly weighted features had marginal impact, which the expert regarded as acceptable but emphasized should be explicit and adjustable.

The expert considered weighted-mean dimension functions mathematically transparent but potentially too rigid for nonlinear relationships, such as threshold effects or super-additive combinations of high data quality and robustness. This suggests exploring alternative aggregation forms (e.g., piecewise-linear or sigmoid functions) in future work. Regarding the scalar norm and its transformation to keep values in $[0, 1]$, the expert found the “harder near 0 and 1” behavior intuitively plausible but noted that, without calibrated decision thresholds, trust scores should be interpreted primarily in a relative, comparative way rather than as absolute guarantees.

5.3. Limitations and Implications

The evaluation is limited to a single expert and simulator-based scenarios, and thus cannot provide statistical validation or fully reflect the messiness of real construction data. Instead, it offers in-depth evidence that the proposed scoring scheme yields trust trajectories that are internally consistent with the defined heuristics and are interpretable for at least one informed user.

The results support the viability of the three-dimensional trust space plus scalar norm as a basis for trust-aware BI, while indicating that (i) careful design and documentation of feature sets and weights are crucial, and (ii) more flexible, context-dependent aggregation functions and norms should be investigated.

6. Conclusion and Outlook

The present work has demonstrated that trustworthiness in AI-based big data analysis can be modeled and implemented in a coherent way by combining the AI2VIS4BigData reference model, the TAI-BDM framework, and a mathematically grounded trust state space as defined in prior work. By instantiating these concepts in the construction domain and embedding them into a fully implemented BI proof-of-concept application, it was possible to make both the evolution of trust along the analysis pipeline and its determinants—trust features, phase-specific aggregations, and an overall trust norm—transparent and inspectable for expert users. The qualitative cognitive walkthrough evaluation indicated that the core ideas of the trust trajectory and the decomposition into trust dimensions are understandable and perceived as useful, while also revealing concrete improvement potentials in GUI design, explanation depth for trust features, and consistency of visual encodings.

The results suggest several directions for future work. First, the current trust dimension functions and overall trust norm are analytically specified and manually calibrated; next steps include learning weights or even functional forms from user behavior, expert feedback, or outcome data (e.g., via Bayesian updating or multi-objective optimization). Second, while the prototype motivates an arithmetic-mean-based trust norm, it neither analyzes conditions under which this choice may fail nor empirically compares alternative norms, so systematic study of such edge cases and aggregation schemes is needed. Third, the trust model has so far been instantiated only for road-construction risk forecasting with a single AI pipeline; validating it across other sectors, model classes, and data regimes, and integrating the trust bus with experiment-tracking and model-management tools (e.g., MLflow-like infrastructures), would help refine feature sets, strengthen FAIR compliance, and support reuse and auditing of models and trust evidence.

The findings from this study are limited to qualitative insights obtained from a single expert in a cognitive walkthrough, and therefore should not be interpreted as broader empirical validation of the proposed trust model. In future work, a multi-user, mixed-method study ought to be conducted that combines controlled quantitative tasks (e.g., trust calibration, error detection, decision quality) with

qualitative techniques across diverse stakeholder groups. This design will also help address threats to the construct validity of ‘trust’, expert selection bias, and potential confirmation bias arising from the authors’ dual role as designers and evaluators.

A promising direction for future research is to use the IELC [17] as a data-ingestion backbone that continuously transforms heterogeneous data, addressing AI’s dependence on high-quality, well-structured data while providing cloud-native orchestration.

Finally, the evaluation in this work was intentionally exploratory and qualitative, with one participant. Scaling this to larger, more diverse user groups and combining qualitative methods with quantitative measures (e.g., task performance, error detection, trust calibration) will be key to assessing how well the proposed trust representations actually support human decision-making in practice. Taken together, these avenues point towards a research program in which the mathematical and architectural foundations of TAI-BDM are iteratively refined and empirically validated across domains, ultimately aiming at a family of trust-aware AI systems where trust trajectories become a first-class object in both design and operation.

Declaration on Generative AI

During the preparation of this work, the authors used ChatGPT and Grammarly for AI-assisted copy editing. ChatGPT was used to provide feedback and suggestions on the clarity and structure of selected sections to improve writing quality. The tool was used solely for suggestions on structure and phrasing, and for critique; no AI-generated text was copied or directly incorporated into the manuscript. All final wording and content are the authors’ own, and the authors take full responsibility for the publication’s content.

References

- [1] T. Reis, M. X. Bornschlegl, M. L. Hemmje, Ai2vis4bigdata: A reference model for ai-based big data analysis and visualization, in: *Advanced Visual Interfaces. Supporting Artificial Intelligence and Big Data Applications: AVI 2020 Workshops, AVI-BDA and ITAVIS*, Ischia, Italy, June 9, 2020 and September 29, 2020, Revised Selected Papers, Springer-Verlag, Berlin, Heidelberg, 2020, p. 1–18. URL: https://doi.org/10.1007/978-3-030-68007-7_1. doi:10.1007/978-3-030-68007-7_1.
- [2] M. X. Bornschlegl, Towards trustworthiness in AI-based big data analysis, 2025. URL: https://ub-deposit.fernuni-hagen.de/receive/mir_mods_00002032, research Program Paper.
- [3] S. Bruchhaus, A. Kreibich, T. Reis, M. X. Bornschlegl, M. L. Hemmje, Towards a conceptual modeling of trustworthiness in AI-based big data analysis, *Information* 16 (2025) 455. URL: <https://doi.org/10.3390/info16060455>. doi:10.3390/info16060455.
- [4] T. Reis, Empowering Users in Visual Big Data Analysis by Applying Artificial Intelligence and Machine Learning, Doctoral dissertation, FernUniversität in Hagen, Hagen, Germany, 2025. URL: https://ub-deposit.fernuni-hagen.de/receive/mir_mods_00002183. doi:10.18445/20250501-103138-0.
- [5] G. Beliakov, A. Pradera, T. Calvo, *Aggregation Functions: A Guide for Practitioners*, Studies in Fuzziness and Soft Computing, Springer, Berlin, 2007.
- [6] A. Kreibich, Vertrauenswürdigkeit von KI-basierten Risikoprognosen in der Baubranche, Master’s thesis, FernUniversität in Hagen, Hagen, Germany, 2026. Supervised by Sebastian Bruchhaus and Matthias L. Hemmje.
- [7] W. Rudin, *Functional Analysis*, International Series in Pure and Applied Mathematics, 2 ed., McGraw-Hill, New York, 1991.
- [8] Python Software Foundation, *Python Language Reference*, 2023. URL: <https://www.python.org/>.
- [9] Spotify, Luigi workflow management system, <https://github.com/spotify/luigi>, 2023.
- [10] M. D. Wilkinson, M. Dumontier, I. J. Aalbersberg, G. Appleton, M. Axton, A. Baak, N. Blomberg, J.-W. Boiten, et al., The FAIR guiding principles for scientific data management and stewardship, *Scientific Data* 3 (2016) 160018. doi:10.1038/sdata.2016.18.

- [11] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, E. Duchesnay, Scikit-learn: Machine learning in Python, *Journal of Machine Learning Research* 12 (2011) 2825–2830.
- [12] S. M. Lundberg, S.-I. Lee, A unified approach to interpreting model predictions, in: *Proceedings of the 31st International Conference on Neural Information Processing Systems, NIPS'17*, Curran Associates Inc., Red Hook, NY, USA, 2017, p. 4768–4777.
- [13] M. T. Ribeiro, S. Singh, C. Guestrin, "why should i trust you?": Explaining the predictions of any classifier, in: *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD '16*, Association for Computing Machinery, New York, NY, USA, 2016, p. 1135–1144. URL: <https://doi.org/10.1145/2939672.2939778>. doi:10.1145/2939672.2939778.
- [14] V. Borisov, T. Leemann, K. Sebler, J. Haug, M. Pawelczyk, G. Kasneci, Deep neural networks and tabular data: A survey, *IEEE Transactions on Neural Networks and Learning Systems* 35 (2024) 7499–7519. doi:10.1109/TNNLS.2022.3229161, publisher Copyright: © 2012 IEEE.
- [15] Vue.js Core Team, Vue.js: The progressive javascript framework, <https://vuejs.org/>, 2023.
- [16] Docker, Inc., Docker: Empowering app development for developers, <https://www.docker.com/>, 2023.
- [17] P. Tamla Mbobda, Towards a reference model for an information extraction lifecycle supporting data collection, data preparation, information extraction, information organization, knowledge organization, information access and retrieval, <https://nbn-resolving.org/urn:nbn:de:hbz:708-dh15437>, 2025. Accessed: 2026-01-21.