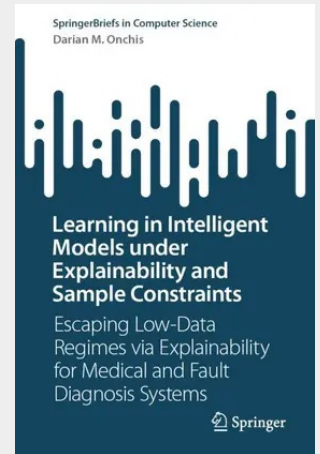


Learning in Intelligent Models under Explainability and Sample Constraints

Escaping Low-Data Regimes via Explainability for Medical and Fault Diagnosis Systems

This book provides a unified framework for learning in intelligent systems under conditions of limited data and strict explainability requirements. It addresses a critical gap in modern machine learning, where high-performance models often rely on large datasets and operate as black boxes, limiting their applicability in high-stakes domains. The book introduces the LIMESC framework, a novel approach that integrates explainability, learning, and domain knowledge into a single methodological structure. It systematically explores how models can remain robust, interpretable, and adaptable when data are scarce, noisy, or evolving. Core topics include neural computing, deep learning under small-data regimes, regularization and optimization strategies, class-incremental learning without memory, and dataset knowledge transfer. The book further examines post-hoc explainability methods and transitions toward intrinsic interpretability through causal and neuro-symbolic approaches. Additional perspectives such as topological data analysis and reinforcement learning in constrained environments are also presented. The framework is grounded in real-world applications, particularly medical diagnostics and fault detection systems, where explainability and reliability are essential. Through a combination of theoretical insights and practical methodologies, the book offers a structured pathway toward designing adaptive and interpretable machine learning systems. This book is intended for researchers, advanced graduate students, and practitioners in machine learning, artificial intelligence, and biomedical engineering.

This book provides a unified framework for learning in intelligent systems under conditions of limited data and strict explainability requirements. It addresses a critical gap in modern machine learning, where high-performance models often rely on large datasets and operate as black boxes, limiting their applicability in high-stakes domains. The book introduces the LIMESC framework, a novel approach that integrates explainability, learning, and domain knowledge into a single methodological structure. It systematically explores how models can remain robust, interpretable, and adaptable when data are scarce, noisy, or evolving. Core topics include neural computing, deep learning under small-data regimes, regularization and optimization strategies, class-incremental learning without memory, and dataset knowledge transfer. The book further examines post-hoc explainability methods and transitions toward intrinsic interpretability through causal and neuro-symbolic approaches. Additional perspectives such as topological data analysis and reinforcement learning in constrained environments are also presented. The framework is grounded in real-world applications, particularly medical diagnostics and fault detection systems, where explainability and reliability are essential. Through a combination of theoretical insights and practical methodologies, the book offers a structured pathway toward designing adaptive and interpretable machine learning systems. This book is intended for researchers, advanced graduate students, and practitioners in machine learning, artificial intelligence, and biomedical engineering.



181,89 €

169,99 € (zzgl. MwSt.)

vorbestellbar, Erscheinungstermin ca. Dezember 2026

Artikelnummer: 9783032408136

Medium: Buch

ISBN: 978-3-032-40813-6

Verlag: Springer

Erscheinungstermin: 20.12.2026

Sprache(n): Englisch

Auflage: Erscheinungsjahr 2026

Serie: SpringerBriefs in Computer Science

Produktform: Kartoniert

Seiten: 130

Format (B x H): 155 x 235 mm

