

benchmarking large language models for automated verilog rtl code generation

benchmarking large language models for automated verilog rtl code generation is an emerging and critical area of research and application within the fields of artificial intelligence and digital design automation. As large language models (LLMs) continue to demonstrate significant capabilities in natural language understanding and code synthesis, their potential to generate hardware description language (HDL) code like Verilog RTL (Register Transfer Level) has garnered substantial interest. This article explores the methodologies, metrics, and challenges involved in benchmarking large language models for automated Verilog RTL code generation. It covers the importance of such benchmarks, the criteria for evaluating model performance, and the practical considerations when integrating LLMs into hardware design workflows. Readers will gain insight into the current state of AI-driven RTL generation and the prospects for improving automation in hardware design. The discussion also addresses semantic accuracy, synthesis readiness, and optimization aspects relevant to Verilog code produced by these models. The following sections provide a comprehensive overview of benchmarking strategies and the implications for future hardware development processes.

- Understanding Large Language Models in Verilog RTL Code Generation
- Key Metrics for Benchmarking Automated Verilog RTL Code Generation
- Benchmarking Methodologies and Frameworks
- Challenges in Benchmarking Large Language Models for RTL Code
- Applications and Future Directions in Automated RTL Code Generation

Understanding Large Language Models in Verilog RTL Code Generation

Large language models, such as GPT and similar transformer-based architectures, have transformed natural language processing and programming automation. In the context of Verilog RTL code generation, these models are trained or fine-tuned to understand hardware design specifications expressed in natural language or domain-specific prompts and translate them into synthesizable Verilog code. This section introduces the key concepts behind LLMs applied to hardware description languages and the significance of automating RTL code generation.

Overview of Large Language Models

Large language models are deep learning models trained on enormous datasets to predict tokens in sequences, enabling them to generate human-like text and code. Their architecture typically involves attention mechanisms that allow for understanding context across long input sequences, which is essential for generating coherent and semantically accurate Verilog RTL code.

Role of LLMs in Hardware Design Automation

Automated Verilog RTL code generation leverages LLMs to reduce the manual effort required in hardware design. By inputting high-level design descriptions or behavioral specifications, LLMs can produce RTL code that adheres to design constraints and logic functionality, accelerating the design cycle and potentially reducing errors.

Key Metrics for Benchmarking Automated Verilog RTL Code Generation

Evaluating the performance of large language models for automated Verilog RTL code generation

necessitates a robust set of metrics that assess code correctness, efficiency, and usability. This section presents the principal metrics used to benchmark these models, facilitating objective comparison and improvement tracking.

Functional Correctness

Functional correctness verifies that the generated Verilog RTL code behaves as intended according to the design specification. Techniques such as simulation against testbenches and formal verification methods are employed to ascertain that the logic implemented by the code matches the expected outcomes.

Code Synthesis and Timing Performance

Beyond correctness, synthesized hardware performance is critical. Benchmarking includes measuring how well the generated RTL code synthesizes in target FPGA or ASIC flows, the timing closure achieved, resource utilization, and power consumption estimates. These factors determine the practical viability of the generated code.

Code Quality and Readability

Code quality metrics evaluate the maintainability and clarity of the Verilog code, including adherence to coding standards, modularity, and absence of redundant or inefficient constructs. Although automated generation focuses on functionality, quality is essential for integration and debugging in real-world projects.

Latency and Throughput of Code Generation

Efficiency of the LLM in producing code is also considered, measuring the inference time, computational resource requirements, and scalability when generating larger or more complex RTL

modules.

Benchmarking Methodologies and Frameworks

Establishing standardized benchmarking methodologies ensures reproducibility and comparability across different large language models designed for Verilog RTL code generation. This section discusses common approaches and frameworks used in the benchmarking process.

Dataset Preparation and Benchmark Suites

Benchmarking requires well-curated datasets consisting of hardware design specifications paired with reference RTL implementations. These datasets may include diverse design patterns, ranging from simple combinational logic to complex sequential circuits, ensuring comprehensive evaluation.

Automated Testing and Verification Pipelines

Integration of automated testing frameworks enables systematic validation of generated RTL code. This includes scripted simulation runs, coverage analysis, and formal equivalence checking to verify that the outputs meet functional and timing requirements.

Comparative Analysis Across Models

Benchmarking frameworks compare multiple LLMs under identical conditions, analyzing differences in generation accuracy, synthesis results, and generation efficiency. This comparative approach helps identify strengths and weaknesses of each model.

Challenges in Benchmarking Large Language Models for RTL Code

Benchmarking automated Verilog RTL code generation presents unique challenges due to the complexity of hardware design and the nuances involved in code synthesis. This section outlines the primary obstacles encountered and their implications.

Semantic Ambiguity in Specifications

Natural language design descriptions can be ambiguous or incomplete, complicating the task of generating precise RTL code. LLMs must interpret specifications accurately, but variations in phrasing or missing details can lead to incorrect or suboptimal code generation.

Evaluation Complexity and Resource Intensity

Comprehensive benchmarking requires extensive simulation and synthesis runs, which can be resource-intensive and time-consuming. Formal verification of generated code also demands significant computational power and expertise.

Balancing Generalization and Specialization

Models trained on general programming corpora may lack domain-specific knowledge essential for high-quality RTL generation. Conversely, highly specialized models might struggle with diverse or novel design requirements, making benchmarking across different design types challenging.

Applications and Future Directions in Automated RTL Code

Generation

Benchmarking large language models for automated Verilog RTL code generation is not only about evaluation but also about guiding future advancements and applications. This section explores current use cases and prospective developments in the field.

Accelerating Hardware Design Cycles

Automated RTL code generation using LLMs can drastically reduce design turnaround times by enabling rapid prototyping and iterative development, supporting agile hardware design methodologies.

Enhancing Design Space Exploration

LLMs can facilitate exploration of multiple design alternatives by generating varied RTL implementations based on different constraints or optimization goals, aiding designers in selecting optimal solutions.

Integration with EDA Tools and Workflows

Future work involves seamless integration of LLM-generated RTL code within electronic design automation (EDA) toolchains, ensuring compatibility with synthesis, place-and-route, and verification tools to streamline end-to-end hardware development.

Advances in Multimodal and Context-Aware Models

Emerging models that combine textual, graphical, and design context inputs promise to improve the accuracy and relevance of automated RTL generation, addressing current limitations in understanding complex hardware specifications.

Key Benefits of Benchmarking Large Language Models for Verilog RTL Code

- Objective assessment of functional and synthesis quality
- Identification of model strengths and areas for improvement
- Promotion of standardized evaluation practices in hardware AI applications
- Facilitation of collaboration between AI researchers and hardware engineers
- Acceleration of innovation through informed model development

Frequently Asked Questions

What is benchmarking in the context of large language models for Verilog RTL code generation?

Benchmarking in this context refers to the systematic evaluation of large language models based on their ability to generate accurate, efficient, and syntactically correct Verilog RTL code, using standardized datasets and performance metrics.

Which metrics are commonly used to benchmark large language models for automated Verilog RTL code generation?

Common metrics include code correctness (functional equivalence), syntax accuracy, code efficiency (resource utilization), inference time, model size, and the ability to handle complex design

specifications.

What are the challenges in benchmarking LLMs for Verilog RTL code generation?

Challenges include the lack of standardized datasets, difficulty in verifying functional correctness, variability in coding styles, the complexity of hardware description languages, and the need for domain-specific evaluation metrics.

How do large language models handle the generation of synthesizable Verilog RTL code?

LLMs are trained on large corpora of Verilog code and learn patterns for generating synthesizable constructs. However, ensuring synthesizability often requires post-generation verification and sometimes human-in-the-loop refinement.

What role do testbenches play in benchmarking automated Verilog RTL code generation?

Testbenches are used to validate the functional correctness of the generated Verilog RTL code by simulating its behavior under various test scenarios, which is critical for benchmarking the model's practical utility.

Are there any publicly available benchmarks or datasets for evaluating Verilog RTL code generation by LLMs?

Currently, publicly available benchmarks are limited, but some research groups have developed datasets comprising Verilog modules and associated specifications, which can be used for training and evaluation purposes.

How does the size and architecture of an LLM impact its performance in generating Verilog RTL code?

Larger models with more parameters generally capture more complex patterns and produce higher-quality code, but they require more computational resources and may have longer inference times, affecting practical deployment.

Can benchmarking help improve the design of LLMs for hardware description language generation?

Yes, benchmarking identifies strengths and weaknesses in model performance, guiding improvements in model architectures, training data quality, and fine-tuning methods tailored to hardware description languages like Verilog.

What future trends are expected in benchmarking large language models for automated Verilog RTL code generation?

Future trends include the development of standardized benchmarks, integration of formal verification techniques, multi-modal models combining code and specification understanding, and real-time code synthesis with feedback loops.

Additional Resources

1. *Benchmarking Large Language Models for Hardware Description Languages*

This book explores the evaluation methodologies for large language models (LLMs) specifically applied to Hardware Description Languages (HDLs) such as Verilog and VHDL. It covers the challenges in measuring code generation accuracy, efficiency, and synthesis readiness. Practical case studies demonstrate benchmarking frameworks and metrics tailored for automated RTL code generation.

2. *Automated RTL Code Generation with Large Language Models*

Focusing on the intersection of AI and hardware design, this book delves into how LLMs can be harnessed to generate RTL code automatically. It discusses model architectures, fine-tuning strategies, and integration into existing hardware design workflows. Readers gain insights into improving code quality and reducing design cycles through automation.

3. Evaluating AI-Driven Verilog Code Synthesis

This title provides an in-depth look at the evaluation criteria and benchmark suites for AI models tasked with generating Verilog code. It highlights key performance indicators such as code correctness, timing closure, and resource utilization. The book also reviews contemporary datasets and challenges in the domain.

4. Large Language Models in Digital Design Automation

A comprehensive guide on leveraging LLMs in digital design automation, this book covers both theoretical foundations and practical applications. It presents methods for integrating AI-generated RTL into EDA tools and discusses benchmarking protocols to assess model effectiveness in real-world scenarios.

5. Towards Reliable Verilog RTL Generation using AI

This publication addresses the reliability and robustness aspects of AI-generated RTL code. It examines error detection, correction mechanisms, and verification techniques to ensure that automatically generated Verilog meets stringent industry standards. Benchmarking approaches to quantify reliability are also discussed.

6. Data-Driven Approaches for Verilog Code Synthesis

Highlighting the role of data in training and benchmarking LLMs, this book investigates dataset creation, augmentation, and annotation for Verilog code generation tasks. It emphasizes the importance of diverse and representative data to improve model generalization and benchmark validity.

7. Performance Metrics for AI-Generated RTL Code

This book focuses entirely on the performance measurement aspects of AI-generated RTL code, proposing novel metrics beyond syntax correctness. It includes discussions on simulation accuracy,

power efficiency, and synthesis feasibility, providing a holistic framework for benchmarking LLM outputs.

8. Integrating Large Language Models into Hardware Design Flows

Offering practical guidance, this book explores how LLMs can be embedded into existing hardware design pipelines to automate RTL code generation. It discusses interoperability challenges, benchmarking integration, and case studies showcasing productivity improvements in design teams.

9. Challenges and Advances in Automated Verilog Code Generation

This book provides an overview of the current challenges in automated Verilog code generation using large language models, including ambiguity in specification interpretation and maintaining design intent. It reviews recent advances in model architectures and benchmarking approaches that address these issues, paving the way for more reliable automation tools.

Benchmarking Large Language Models For Automated Verilog Rtl Code Generation

Find other PDF articles:

<https://admin.nordenson.com/archive-library-406/files?trackid=LhD52-2277&title=if-you-give-a-mouse-a-cookie-teacher-gift.pdf>

benchmarking large language models for automated verilog rtl code generation: Next Generation EDA Flow Khaled Salah Mohamed, 2025-05-13 This book serves as a comprehensive guide to the world of EDA tools, offering readers a deeper understanding of their inner workings and a glimpse into the future of electronic design. With a meticulous focus on numerical methods, the author delves deeply into the mathematical foundations that underpin EDA tools. From finite element analysis to Monte Carlo simulations, readers will gain a thorough understanding of the numerical techniques employed to model and simulate complex electronic systems. Furthermore, this book elucidates the diverse modeling methods utilized in EDA tools, providing readers with a holistic view of the methods employed to represent and analyze electronic circuits and systems. Whether exploring circuit-level simulations or system-level modeling, readers will be equipped with the knowledge needed to navigate the intricacies of EDA toolsets. The author also delves into the fascinating intersection of quantum mechanics and electronic design, examining the evolving landscape of quantum EDA tools and offering insights into the transformative potential of quantum computing in electronic design. Lastly, this book explores the transformative impact of machine learning on EDA tools, offering insights into how artificial intelligence techniques can enhance performance and productivity.

benchmarking large language models for automated verilog rtl code generation: Applied Cryptography and Network Security Workshops Martin Andreoni, 2024-06-23 This two-volume set LNCS 14586-14587 constitutes the proceedings of eight Satellite Workshops held in parallel with the 22nd International Conference on Applied Cryptography and Network Security, ACNS 2024, held in Abhu Dabhi, United Arab Emirates, during March 5-8, 2024. The 33 full papers and 11 poster papers presented in this volume were carefully reviewed and selected from 62 submissions. They stem from the following workshops: 6th ACNS Workshop on Application Intelligence and Blockchain Security (AIBlock 2024). 5th ACNS Workshop on Artificial Intelligence in Hardware Security (AIHWS 2024). 6th ACNS Workshop on Artificial Intelligence and Industrial IoT Security (AIoTS 2024). 5th ACNS Workshop on Secure Cryptographic Implementation (SCI 2024). 1st Workshop on Advances in Asymmetric Cryptanalysis (AAC 2024). 6th ACNS Workshop on Security in Machine Learning and its Applications (SiMLA 2024). 1st Workshop on Low-Latency Encryption (LLE 2024). 4th ACNS Workshop on Critical Infrastructure and Manufacturing System Security (CIMSS 2024).

benchmarking large language models for automated verilog rtl code generation: **Bridging the Gap Between AI and Reality** Bernhard Steffen, 2025-09-30 This open access book constitutes revised selected papers from the Second International Conference on Bridging the Gap between AI and Reality, AISoLA 2024, which took place in Crete, Greece, in October/November 2024. The papers included in this book extend the presentation in the AISoLA 2024 on-site proceedings. They focus on the following topics: AI-Assisted Programming; health care approaches using formal methods and AI; responsible and trusted AI: an interdisciplinary perspective; statistical model checking; and verification for neur-symbolic artificial intelligence.

benchmarking large language models for automated verilog rtl code generation: AI-Enabled Electronic Circuit and System Design Ali Iranmanesh, Hossein Sayadi, 2025-01-27 As our world becomes increasingly digital, electronics underpin nearly every industry. Understanding how AI enhances this foundational technology can unlock innovations, from smarter homes to more powerful gadgets, offering vast opportunities for businesses and consumers alike. This book demystifies how AI streamlines the creation of electronic systems, making them smarter and more efficient. With AI's transformative impact on various engineering fields, this resource provides an up-to-date exploration of these advancements, authored by experts actively engaged in this dynamic field. Stay ahead in the rapidly evolving landscape of AI in engineering with "AI-Enabled Electronic Circuit and System Design: From Ideation to Utilization," your essential guide to the future of electronic systems. !--[endif]--A transformative guide describing how revolutionizes electronic design through AI integration. Highlighting trends, challenges and opportunities; Demystifies complex AI applications in electronic design for practical use; Leading insights, authored by top experts actively engaged in the field; Offers a current, relevant exploration of significant topics in AI's role in electronic circuit and system design. Editor's bios. Dr. Ali A. Iranmanesh is the founder and CEO of Silicon Valley Polytechnic Institute. He has received his Bachelor of Science in Electrical Engineering from Sharif University of Technology (SUT), Tehran, Iran, and both his master's and Ph.D. degrees in Electrical Engineering and Physics from Stanford University in Stanford, CA. He additionally holds a master's degree in business administration (MBA) from San Jose State University in San Jose, CA. Dr. Iranmanesh is the founder and chairman of the International Society for Quality Electronic Design (ISQED). Currently, he serves as the CEO of Innovotek. Dr. Iranmanesh has been instrumental in advancing semiconductor technologies, innovative design methodologies, and engineering education. He holds nearly 100 US and international patents, reflecting his significant contributions to the field. Dr. Iranmanesh is the Senior life members of IEEE, senior member of the American Society for Quality, co-founder and Chair Emeritus of the IEEE Education Society of Silicon Valley, Vice Chair Emeritus of the IEEE PV chapter, and recipient of IEEE Outstanding Educator Award. Dr. Hossein Sayadi is a Tenure-Track Assistant Professor and Associate Chair in the Department of Computer Engineering and Computer Science at California State University, Long Beach (CSULB). He earned his Ph.D. in Electrical and Computer Engineering from George Mason University in Fairfax, Virginia, and an M.Sc. in Computer Engineering from Sharif University of

Technology in Tehran, Iran. As a recognized researcher with over 14 years of research experience, Dr. Sayadi is the founder and director of the Intelligent, Secure, and Energy-Efficient Computing (iSEC) Lab at CSULB. His research focuses on advancing hardware security and trust, AI and machine learning, cybersecurity, and energy-efficient computing, addressing critical challenges in modern computing and cyber-physical systems. He has authored over 75 peer-reviewed publications in leading conferences and journals. Dr. Sayadi is the CSU STEM-NET Faculty Fellow, with his research supported by multiple National Science Foundation (NSF) grants and awards from CSULB and the CSU Chancellor's Office. He has contributed to various international conferences as an organizer and program committee member, including as the TPC Chair for the 2024 and 2025 IEEE ISQED.

benchmarking large language models for automated verilog rtl code generation:

Digitaltechnik und Digitale Systeme Jürgen Reichardt, 2025-09-22 Mit diesem, nun in der 6. Auflage vorliegenden Buch wird erstmalig ein Bogen von den Grundlagen der Digitaltechnik über den VHDL-basierten Schaltungsentwurf bis zur High-Level-Synthese aufgespannt. Es stellt somit ein aktuelles Kompendium zu rechnergestützten Methoden beim Entwurf digitaler Systeme dar. Neben dem Verständnis des funktionellen und zeitlichen Verhaltens von Logikfunktionen und grundlegenden Implementierungskonzepten für digitale Systeme wird auch die Fertigkeit zur synthesesgerechten Modellierung mit einer Hardwarebeschreibungssprache wie VHDL oder einer Programmiersprache wie C bzw. C++ vermittelt. Mit der signifikanten Ergänzung zur High-Level Synthese wird der Tatsache Rechnung getragen, dass sich die Ergebnisqualität der zugehörigen Synthesecompiler in den letzten Jahren deutlich verbessert hat, sodass auf C/C++-Ebene beschriebene Systeme durch automatisch vorgenommene Optimierungen einen erheblichen Performancegewinn bringen, der bei Implementierungen mit Hardwarebeschreibungssprachen noch manuell hinzugefügt werden muss. Der ausgezeichnete didaktische Aufbau unterstützt den Lernprozess: Den Kapiteln sind Lernziele vorangestellt und immer wieder werden grafische und tabellarische Übersichten sowie vertiefende Beispiele präsentiert. Eine Vielzahl von Übungsaufgaben mit Musterlösungen dient zur Lernkontrolle.

Related to benchmarking large language models for automated verilog rtl code generation

Benchmarking Large Language Models for Automated Verilog In this section, we describe our method for training (or fine-tuning) LLM models for Verilog code generation. We begin by describing our curated Verilog datasets, followed by the LLM

Benchmarking Large Language Models for Automated Verilog RTL Code Automating hardware design could obviate a significant amount of human error from the engineering process and lead to fewer errors. Verilog is a popular hardware

GitHub - shailja-thakur/VGen In this paper, we characterize the ability of LLMs to generate useful Verilog. For this, we fine-tune pre-trained LLMs on Verilog datasets collected from GitHub and Verilog textbooks

Benchmarking Large Language Models for Code Generation This research endeavors to explore the effectiveness of Large Language Models (LLMs) in code generation, focusing on benchmarking their performance across various coding tasks

VeriGen: A Large Language Model for Verilog Code Generation In this study, we explore the capability of Large Language Models (LLMs) to automate hardware design by automatically completing partial Verilog code, a common

Benchmarking Large Language Models for Automated Verilog e language models (LLMs) are able to write high-quality code in other programming languages. In this paper, we characterize the ability of LLMs to generate useful Verilog. For this, we fine-tune

Benchmarking Large Language Models for Automated Verilog RTL Code Emerging large language models (LLMs) are able to write high-quality code in other programming languages. In this

paper, we characterize the ability of LLMs to generate useful Verilog

VerilogEval: Evaluating Large Language Models for Verilog Code Generation The

increasing popularity of large language models (LLMs) has paved the way for their application in diverse domains. This paper proposes a benchmarking framework tailored

VeriBench: Benchmarking Large Language Models for Verilog Code In the rapidly advancing field of hardware design, Electronic Design Automation (EDA) tools can be significantly improved using Machine Learning. This study eva

Natural language is not enough: Benchmarking multi-modal Abstract Natural language interfaces have exhibited considerable potential in the automation of Verilog generation derived from high-level specifications through the utilization

Benchmarking Large Language Models for Automated In this section, we describe our method for training (or fine-tuning) LLM models for Verilog code generation. We begin by describing our curated Verilog datasets, followed by the LLM

Benchmarking Large Language Models for Automated Verilog RTL Code Automating hardware design could obviate a significant amount of human error from the engineering process and lead to fewer errors. Verilog is a popular hardwa

GitHub - shailja-thakur/VGen In this paper, we characterize the ability of LLMs to generate useful Verilog. For this, we fine-tune pre-trained LLMs on Verilog datasets collected from GitHub and Verilog textbooks

Benchmarking Large Language Models for Code Generation This research endeavors to explore the effectiveness of Large Language Models (LLMs) in code generation, focusing on benchmarking their performance across various coding tasks

VeriGen: A Large Language Model for Verilog Code Generation In this study, we explore the capability of Large Language Models (LLMs) to automate hardware design by automatically completing partial Verilog code, a common

Benchmarking Large Language Models for Automated e language models (LLMs) are able to write high-quality code in other programming languages. In this pap r, we characterize the ability of LLMs to generate useful Verilog. For this, we fine-tune

Benchmarking Large Language Models for Automated Verilog RTL Code Emerging large language models (LLMs) are able to write high-quality code in other programming languages. In this paper, we characterize the ability of LLMs to generate useful Verilog

VerilogEval: Evaluating Large Language Models for Verilog Code Generation The increasing popularity of large language models (LLMs) has paved the way for their application in diverse domains. This paper proposes a benchmarking framework tailored

VeriBench: Benchmarking Large Language Models for Verilog Code In the rapidly advancing field of hardware design, Electronic Design Automation (EDA) tools can be significantly improved using Machine Learning. This study eva

Natural language is not enough: Benchmarking multi-modal Abstract Natural language interfaces have exhibited considerable potential in the automation of Verilog generation derived from high-level specifications through the utilization of

Benchmarking Large Language Models for Automated Verilog In this section, we describe our method for training (or fine-tuning) LLM models for Verilog code generation. We begin by describing our curated Verilog datasets, followed by the LLM

Benchmarking Large Language Models for Automated Verilog RTL Code Automating hardware design could obviate a significant amount of human error from the engineering process and lead to fewer errors. Verilog is a popular hardwa

GitHub - shailja-thakur/VGen In this paper, we characterize the ability of LLMs to generate useful Verilog. For this, we fine-tune pre-trained LLMs on Verilog datasets collected from GitHub and Verilog textbooks

Benchmarking Large Language Models for Code Generation This research endeavors to explore the effectiveness of Large Language Models (LLMs) in code generation, focusing on

benchmarking their performance across various coding tasks

VeriGen: A Large Language Model for Verilog Code Generation In this study, we explore the capability of Large Language Models (LLMs) to automate hardware design by automatically completing partial Verilog code, a common

Benchmarking Large Language Models for Automated Verilog e language models (LLMs) are able to write high-quality code in other programming languages. In this paper, we characterize the ability of LLMs to generate useful Verilog. For this, we fine-tune

Benchmarking Large Language Models for Automated Verilog RTL Code Emerging large language models (LLMs) are able to write high-quality code in other programming languages. In this paper, we characterize the ability of LLMs to generate useful Verilog

VerilogEval: Evaluating Large Language Models for Verilog Code Generation The increasing popularity of large language models (LLMs) has paved the way for their application in diverse domains. This paper proposes a benchmarking framework tailored

VeriBench: Benchmarking Large Language Models for Verilog Code In the rapidly advancing field of hardware design, Electronic Design Automation (EDA) tools can be significantly improved using Machine Learning. This study eval

Natural language is not enough: Benchmarking multi-modal Abstract Natural language interfaces have exhibited considerable potential in the automation of Verilog generation derived from high-level specifications through the utilization

Back to Home: <https://admin.nordenson.com>