



Intel Releases Open Source AI Reference Kits

Open source designs simplify AI development for solutions across healthcare, manufacturing, retail and other industries.

SANTA CLARA, Calif.--(BUSINESS WIRE)-- **What's New:** Intel has released the first set of open source AI reference kits specifically designed to make AI more accessible to organizations in on-prem, cloud and edge environments. First introduced at [Intel Vision](#), the reference kits include AI model code, end-to-end machine learning pipeline instructions, libraries and Intel oneAPI components for cross-architecture performance. These kits enable data scientists and developers to learn how to deploy AI faster and more easily across healthcare, manufacturing, retail and other industries with higher accuracy, better performance and lower total cost of implementation.

"Innovation thrives in an open, democratized environment. The Intel accelerated open AI software ecosystem including optimized popular frameworks and Intel's AI tools are built on the foundation of an open, standards-based, unified oneAPI programming model. These reference kits, built with components of Intel's end-to-end AI software portfolio, will enable millions of developers and data scientists to introduce AI quickly and easily into their applications or boost their existing intelligent solutions."

—Wei Li, Ph.D., Intel vice president and general manager of AI and Analytics

About AI Reference Kits: AI workloads continue to grow and diversify with use cases in vision, speech, recommender systems and more. Intel's AI reference kits, built in collaboration with Accenture, are designed to accelerate the adoption of AI across industries. They are open source, pre-built AI with meaningful enterprise contexts for both greenfield AI introduction and strategic changes to existing AI solutions.

Four kits are available for download today:

- **[Utility asset health](#):** As energy consumption continues to grow worldwide, power distribution assets in the field are expected to grow. This predictive analytics model was trained to help utilities deliver higher service reliability. It uses Intel-optimized XGBoost through the [Intel® oneAPI Data Analytics Library](#) to model the health of utility poles with 34 attributes and more than 10 million data points¹. Data includes asset age, mechanical properties, geospatial data, inspections, manufacturer, prior repair and maintenance history, and outage records. The predictive asset maintenance model continuously learns as new data, like new pole manufacturer, outages and other changes in condition, are provided.
- **[Visual quality control](#):** Quality control (QC) is essential in any manufacturing operation. The challenge with computer vision techniques is that they often require heavy graphics compute power during training and frequent retraining as new products

are introduced. The AI Visual QC model was trained using [Intel® AI Analytics Toolkit](#), including Intel® Optimization for PyTorch and [Intel® Distribution of OpenVINO™ toolkit](#), both powered by oneAPI to optimize training and inferencing to be 20% and 55% faster, respectively, compared to stock implementation of Accenture visual quality control kit without Intel optimizations² for computer vision workloads across CPU, GPU and other accelerator-based architectures. Using computer vision and SqueezeNet classification, the AI Visual QC model used hyperparameter tuning and optimization to detect pharmaceutical pill defects with 95% accuracy.

- **Customer chatbot:** Conversational chatbots have become a critical service to support initiatives across the enterprise. AI models that support conversational chatbot interactions are massive and highly complex. This reference kit includes deep learning natural language processing models for intent classification and named-entity recognition using BERT and PyTorch. [Intel® Extension for PyTorch](#) and Intel Distribution of OpenVINO toolkit optimize the model for better performance – 45% faster inferencing compared to stock implementation of Accenture customer chatbot kit without Intel optimizations³ – across heterogeneous architectures, and allow developers to reuse model development code with minimal code changes for training and inferencing.
- **Intelligent document indexing:** Enterprises process and analyze millions of documents every year, and many of the semi-structured and unstructured documents are routed manually. AI can automate the processing and categorizing of these documents for faster routing and lower manual labor costs. Using a support vector classification (SVC) model, this kit was optimized with [Intel® Distribution of Modin](#) and [Intel® Extension for Scikit-learn](#) powered by oneAPI. These tools improve data pre-processing, training and inferencing times to be 46%, 96% and 60% faster, respectively, compared to stock implementation of Accenture Intelligent document indexing kit without Intel optimizations⁴ for reviewing and sorting the documents at 65% accuracy.

Download free on the [Intel.com AI Reference Kits website](#). The kits are also available on Github.

Why It Matters: Developers are looking to infuse AI into their solutions and the reference kits contribute to that goal. These kits build on and complement Intel's AI software portfolio of end-to-end tools and framework optimizations. Built on the foundation of the [oneAPI](#) open, standards-based, heterogeneous programming model, which delivers performance across multiple types of architectures, these tools help data scientists train models faster and at lower cost by overcoming the limitations of proprietary environments.

What's Next: Over the next year, Intel will release a series of additional open source AI reference kits with trained machine learning and deep learning models to help organizations of all sizes in their digital transformation journey.

More Context: [oneAPI Dev Summit for AI](#) | [Intel oneAPI](#) | [Intel AI Tools](#)

About Intel

Intel (Nasdaq: INTC) is an industry leader, creating world-changing technology that enables global progress and enriches lives. Inspired by Moore's Law, we continuously work to

advance the design and manufacturing of semiconductors to help address our customers' greatest challenges. By embedding intelligence in the cloud, network, edge and every kind of computing device, we unleash the potential of data to transform business and society for the better. To learn more about Intel's innovations, go to newsroom.intel.com and intel.com.

Notices & Disclaimers

¹Predictive Utility Analytics Reference Kit, measured on June 29, 2022. HW Configuration: Microsoft Azure Standard D4_v5, OS: Ubuntu 20.04.4 LTS (Focal Fossa), 8 X Intel® Xeon® Platinum 8370C CPU @ 2.80GHz, 2 threads/core, 4 cores/socket, 1 socket. SW Configuration: Config 1 (Python v3.9, Scikit-learn v 1.0.2, XGBoost v0.81), Config 2 (Intel® Distribution for Python 3.9.12 2022.0.0, Scikit-learn 0.24.2, Intel® Extension for Scikit-learn 2021.5.1, XGBoost 1.4.3, daap4py 2021.6.0). Additional details at <https://github.com/oneapi-src/predictive-health-analytics>. Results may vary.

²Visual Quality Inspection Reference Kit, measured on June 29, 2022. HW Configuration: Microsoft Azure Standard D4_v5, OS: Ubuntu 20.04.4 LTS (Focal Fossa), 4 X Intel® Xeon® Platinum 8370C CPU @ 2.80GHz, 2 threads/core, 2 cores/socket, 1 socket. SW Configuration: Config 1 (PyTorch v1.8.0), Config 2 (Intel® Extension for PyTorch v1.8.0, Intel® Neural Compressor v1.12, Intel® Distribution of OpenVINO Toolkit 2021.4.2). Additional details at <https://github.com/oneapi-src/visual-quality-inspection>. Results may vary.

³Customer Chatbot Reference Kit, measured on June 22, 2022. HW Configuration: Microsoft Azure Standard D4_v5, OS: Red Hat Enterprise Linux Server 7.9, 4 X Intel® Xeon® Platinum 8370C CPU @ 2.80GHz, 2 threads/core, 2 cores/socket, 1 socket. SW Configuration: Config 1 (PyTorch v1.11), Config 2 (PyTorch v1.11.0, Intel® Extension for PyTorch v1.11.200, Intel® Neural Compressor v1.12). Additional details at <https://github.com/oneapi-src/customer-chatbot>. Results may vary.

⁴Intelligent Indexing Reference Kit, measured on June 22, 2022. HW Configuration: Amazon AWS m6i.xlarge, OS: Red Hat Enterprise Linux Server 7.9, 4 X Intel® Xeon® Platinum 8370C CPU @ 2.80GHz, 2 threads/core, 2 cores/socket, 1 socket. SW Configuration: Config 1 (Pandas, Scikit-learn), Config 2 (Intel® AI Analytics Toolkit v 2021.4.1, Intel® Extension for Scikit-learn, Intel® Distribution of Modin). Additional details at <https://github.com/oneapi-src/intelligent-indexing>. Results may vary.

Performance varies by use, configuration and other factors. Learn more at [www.Intel.com/PerformanceIndex](https://www.intel.com/PerformanceIndex).

Results may vary. Performance results are based on testing as of dates shown in configurations and may not reflect all publicly available updates.

No product or component can be absolutely secure.

Your costs and results may vary.

Intel technologies may require enabled hardware, software or service activation.

Intel does not control or audit third-party data. You should consult other sources to evaluate

accuracy.

© Intel Corporation. Intel, the Intel logo and other Intel marks are trademarks of Intel Corporation or its subsidiaries. Other names and brands may be claimed as the property of others.

View source version on businesswire.com:

<https://www.businesswire.com/news/home/20220712005386/en/>

Rebecca Want

1-650-919-4595

becky.want@ketchum.com

Source: Intel Corporation