

March 21, 2018



Intel Editorial: One Simple Truth about Artificial Intelligence in Healthcare: It's Already Here

Recent Progress with Intel-Powered Solutions Shows Artificial Intelligence Is Good Medicine.

SAN FRANCISCO--(BUSINESS WIRE)-- The following is an opinion editorial provided by Navin Shenoy, executive vice president and general manager of the Data Center Group at Intel Corporation.

This press release features multimedia. View the full release here:

<https://www.businesswire.com/news/home/20180321006119/en/>

Navin Shenoy, executive vice president and general manager of the Data Center Group at Intel Corporation, speaks Wednesday, March 21, 2018, at the University of California, San Francisco. Intel invited leading healthcare organizations to address the most pressing topics and challenges in healthcare. (Credit: Intel Corporation)

In the wide world of big data, artificial intelligence (AI) holds transformational promise. Everything from manufacturing to

transportation to retail to education will be improved through its application. But nowhere is that potential more profound than in [healthcare](#), where every one of us has a stake.

What if we could predict the next big disease epidemic, and stop it before it kills? What if we could look at zettabytes of data to find those at greatest risk of becoming sick, then quickly and precisely prevent that from happening? What if the treatment and management of chronic disease could be so personalized that no two individuals get the same medicine, but equally enjoy the best possible outcome? What if we could drastically reduce the time and cost to discover new drugs and bring them on the market? What if we could do all of that now?

Thanks to [artificial intelligence](#) and the work of [Intel and its partners](#), we can.

Real Impact Today

There's a common myth that AI in healthcare is the stuff of science fiction – think machines diagnosing illness and prescribing treatment without a doctor involved. But that is not only highly unlikely, it's not even close to the best examples of how AI is emerging in healthcare today.

Intel and partners throughout the healthcare industry – including [GE Healthcare](#), [Siemens](#), [Sharp Healthcare](#), the [Broad Institute](#), [UCSF](#) and the [Mayo Clinic](#) – are successfully applying AI solutions today, from the back office to the doctor's office, from the emergency room to the living room. A few customers that we're working closely with include:

Montefiore Medical System: using prescriptive models to identify patients at risk for respiratory failure, so healthcare workers can act on alerts that lead to timely interventions that save lives and resources.

Stanford Medical: using AI to augment MRI image reconstruction so that a complete image can be delivered in about a minute versus what normally would take about an hour – eliminating risky intubation and sedation in pediatric patients during imaging exams.

ICON plc: Instead of only relying on burdensome clinic visits and paper diaries, using clinical data from sensors and wearable devices to more quickly assess the impact of new therapies in clinical trials.

AccuHealth: Using home monitoring along with data mining and predictive modeling to identify changes of concern among chronic disease patients to enable intervention before conditions escalate and become acute.

Better Health Tomorrow

But the triumph of artificial intelligence in healthcare isn't inevitable. Right now, the average hospital generates 665 terabytes of data annually,¹ but most of that data isn't useful. At least 80 percent of hospital data is unstructured,² such as clinical notes, video and images. Electronic medical records (EMRs) are a mandated system of record, but they aren't as actionable as they could be. Only with AI can we leverage healthcare data to create a system of insights.

Getting healthcare systems to provide greater access to their data would help. Government also has a role to play by providing appropriate incentives and regulatory clarity for sharing data. We agree with the recent [White House proposal](#) to give patients control and ownership of all their health data, bringing it with them wherever they go rather than residing in various doctor's offices, clinics and hospitals.

New technology can help as well. One example: Intel researchers are making great strides toward practical methods for homomorphic encryption, a method that will allow computer systems to perform calculations on encrypted information without decrypting it first. Such encryption would enable researchers to operate on data in a secure and private way, while still delivering insightful results.

Indeed, much work is ahead, and Intel is [uniquely positioned](#) to help healthcare organizations succeed. Emerging healthcare data is massive data – images, a growing list of 'omics (i.e. genomics, proteomics), video – and will require a storage plan and a network that addresses speed, latency and reliability. We have been investing with our partners to build the right systems - data, storage, network, full infrastructure – all the way from the edge to the network to the cloud, and everywhere in between. With the advancements in our hardware and optimizations of popular deep learning frameworks, the Intel Xeon Scalable processor has 198x better inference performance and 127x better training performance than prior generations³. As a result, the Xeon platform is at the center of many AI workloads that are real today because it is well suited for many machine and deep learning applications across industries like healthcare.

But hardware, storage and network alone are not enough. We need to leverage the

unparalleled expertise from data scientists, software developers, industry experts, and ecosystem partners —to address AI in healthcare end to end. As part of the effort to expand expertise across AI, we launched the Intel AI Academy, a place that offers learning materials, community tools and technology to boost AI developments. With more than 250K monthly participants, I invite you to [join for free as well](#).

I feel very fortunate to work for a company like Intel that is committed to powering AI solutions that will tackle some of the biggest challenges of our time, including healthcare. I'm also proud to be leading the team that will deliver that vision.

¹ Source: <http://www.netapp.com/us/media/wp-7169.pdf>

² Source: <http://www.zdnet.com/news/unstructured-data-challenge-or-asset/6356681>

³ Source: Configuration: AI Performance – Software + Hardware (see below)

- INFERENCE using FP32 Batch Size Caffe GoogleNet v1 256 AlexNet 256.
- Performance estimates were obtained prior to implementation of recent software patches and firmware updates intended to address exploits referred to as "Spectre" and "Meltdown." Implementation of these updates may make these results inapplicable to your device or system. Software and workloads used in performance tests may have been optimized for performance only on Intel microprocessors. Performance tests, such as SYSmark and MobileMark, are measured using specific computer systems, components, software, operations and functions. Any change to any of those factors may cause the results to vary. You should consult other information and performance tests to assist you in fully evaluating your contemplated purchases, including the performance of that product when combined with other products. For more complete information visit <http://www.intel.com/performance>
Source: Intel measured as of June 2017 Optimization Notice: Intel's compilers may or may not optimize to the same degree for non-Intel microprocessors for optimizations that are not unique to Intel microprocessors. These optimizations include SSE2, SSE3, and SSSE3 instruction sets and other optimizations. Intel does not guarantee the availability, functionality, or effectiveness of any optimization on microprocessors not manufactured by Intel. Microprocessor-dependent optimizations in this product are intended for use with Intel microprocessors. Certain optimizations not specific to Intel microarchitecture are reserved for Intel microprocessors. Please refer to the applicable product User and Reference Guides for more information regarding the specific instruction sets covered by this notice.
- **Configurations for Inference throughput**
- Processor :2 socket Intel(R) Xeon(R) Platinum 8180 CPU @ 2.50GHz / 28 cores HT ON , Turbo ON Total Memory 376.46GB (12slots / 32 GB / 2666 MHz).CentOS Linux-7.3.1611-Core , SSD sda RS3WC080 HDD 744.1GB,sdb RS3WC080 HDD 1.5TB,sdc RS3WC080 HDD 5.5TB , Deep Learning Framework caffe
version: f6d01efbe93f70726ea3796a4b89c612365a6341 Topology :googlenet_v1
BIOS:SE5C620.86B.00.01.0004.071220170215 MKLDNN:
version: ae00102be506ed0fe2099c6557df2aa88ad57ec1 NoDataLayer. Measured:

1190 imgs/sec vs Platform: 2S Intel® Xeon® CPU E5-2699 v3 @ 2.30GHz (18 cores), HT enabled, turbo disabled, scaling governor set to “performance” via intel_pstate driver, 256GB DDR4-2133 ECC RAM. CentOS Linux release 7.3.1611 (Core), Linux kernel 3.10.0-514.el7.x86_64. OS drive: Seagate* Enterprise ST2000NX0253 2 TB 2.5" Internal Hard Drive. Performance measured with: Environment variables: KMP_AFFINITY='granularity=fine, compact,1,0', OMP_NUM_THREADS=36, CPU Freq set with cpupower frequency-set -d 2.3G -u 2.3G -g performance. Deep Learning Frameworks: Intel Caffe: (<http://github.com/intel/caffe/>), revision b0ef3236528a2c7d2988f249d347d5fdae831236. Inference measured with “caffe time -forward_only” command, training measured with “caffe time” command. For “ConvNet” topologies, dummy dataset was used. For other topologies, data was stored on local storage and cached in memory before training. Topology specs from https://github.com/intel/caffe/tree/master/models/intel_optimized_models (GoogLeNet, AlexNet, and ResNet-50), https://github.com/intel/caffe/tree/master/models/default_vgg_19 (VGG-19), and https://github.com/soumith/convnet-benchmarks/tree/master/caffe/imagenet_winners (ConvNet benchmarks; files were updated to use newer Caffe prototxt format but are functionally equivalent). GCC 4.8.5, MKLML version 2017.0.2.20170110. BVLC-Caffe: <https://github.com/BVLC/caffe>, Inference & Training measured with “caffe time” command. For “ConvNet” topologies, dummy dataset was used. For other topologies, data was stored on local storage and cached in memory before training BVLC Caffe (<http://github.com/BVLC/caffe>), revision 91b09280f5233caf62954c98ce8bc4c204e7475 (commit date 5/14/2017). BLAS: atlas ver. 3.10.1.

- **Configuration for training throughput:**

- Processor :2 socket Intel(R) Xeon(R) Platinum 8180 CPU @ 2.50GHz / 28 cores HT ON , Turbo ON Total Memory 376.28GB (12slots / 32 GB / 2666 MHz).CentOS Linux-7.3.1611-Core , SSD sda RS3WC080 HDD 744.1GB,sdb RS3WC080 HDD 1.5TB,sdc RS3WC080 HDD 5.5TB , Deep Learning Framework caffe version: f6d01efbe93f70726ea3796a4b89c612365a6341 Topology :alexnet BIOS:SE5C620.86B.00.01.0009.101920170742 MKLDNN: version: ae00102be506ed0fe2099c6557df2aa88ad57ec1 NoDataLayer. Measured: 1023 imgs/sec vs Platform: 2S Intel® Xeon® CPU E5-2699 v3 @ 2.30GHz (18 cores), HT enabled, turbo disabled, scaling governor set to “performance” via intel_pstate driver, 256GB DDR4-2133 ECC RAM. CentOS Linux release 7.3.1611 (Core), Linux kernel 3.10.0-514.el7.x86_64. OS drive: Seagate* Enterprise ST2000NX0253 2 TB 2.5" Internal Hard Drive. Performance measured with: Environment variables: KMP_AFFINITY='granularity=fine, compact,1,0', OMP_NUM_THREADS=36, CPU Freq set with cpupower frequency-set -d 2.3G -u 2.3G -g performance. Deep Learning Frameworks: Intel Caffe: (<http://github.com/intel/caffe/>), revision b0ef3236528a2c7d2988f249d347d5fdae831236. Inference measured with “caffe time -forward_only” command, training measured with “caffe time” command. For “ConvNet” topologies, dummy dataset was used. For other topologies, data was stored on local storage and cached in memory before training. Topology specs from https://github.com/intel/caffe/tree/master/models/intel_optimized_models (GoogLeNet, AlexNet, and ResNet-50), https://github.com/intel/caffe/tree/master/models/default_vgg_19 (VGG-19), and https://github.com/soumith/convnet-benchmarks/tree/master/caffe/imagenet_winners

(ConvNet benchmarks; files were updated to use newer Caffe prototxt format but are functionally equivalent). GCC 4.8.5, MKLML version 2017.0.2.20170110. BVLC-Caffe: <https://github.com/BVLC/caffe>, Inference & Training measured with “caffe time” command. For “ConvNet” topologies, dummy dataset was used. For other topologies, data was stored on local storage and cached in memory before training BVLC Caffe (<http://github.com/BVLC/caffe>), revision 91b09280f5233cafc62954c98ce8bc4c204e7475 (commit date 5/14/2017). BLAS: atlas ver. 3.10.1.)

*Other names and brands may be claimed as the property of others.

View source version on businesswire.com:

<https://www.businesswire.com/news/home/20180321006119/en/>

Intel Corporation

Robin Holt

Robin.holt@intel.com

Source: Intel Corporation