



life.augmented

Introducing ST Neural-ART Accelerator

Enabling high-end, power-efficient edge AI
performance on MCUs.

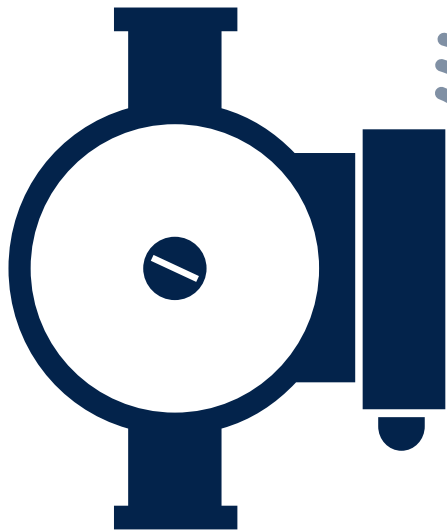
Redefining product greatness



AI drives smarter products and new business models

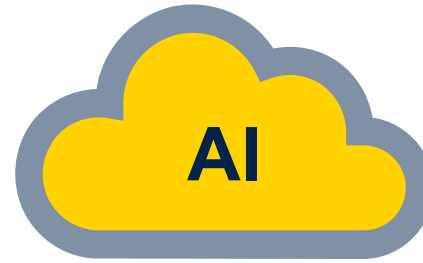
Cloud processing for AI & IoT: Generating a tsunami of data

**Cloud based AI
(IoT devices)**



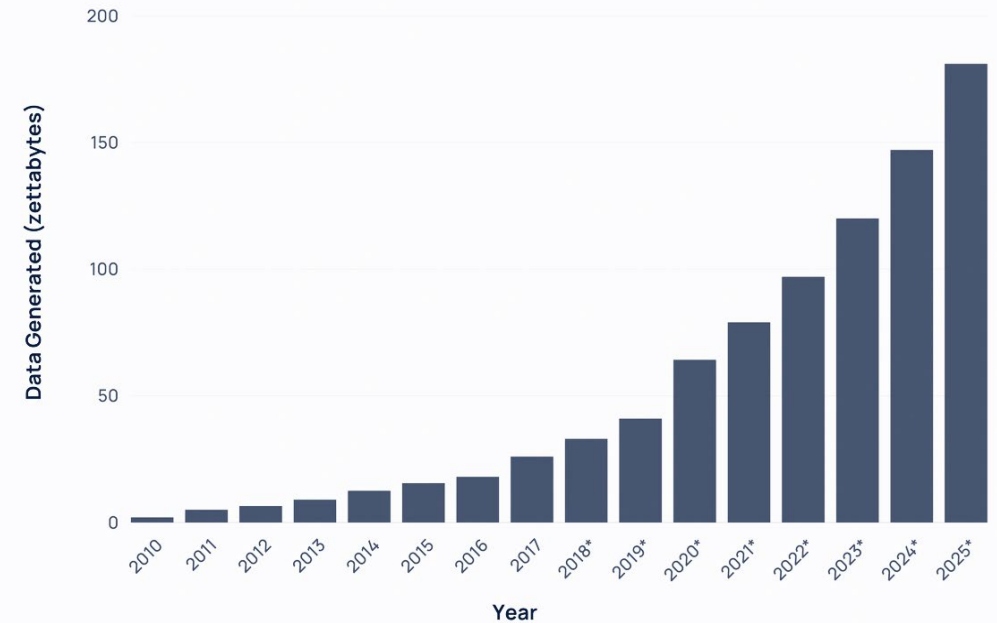
Raw data

Results



AI inference
& storage

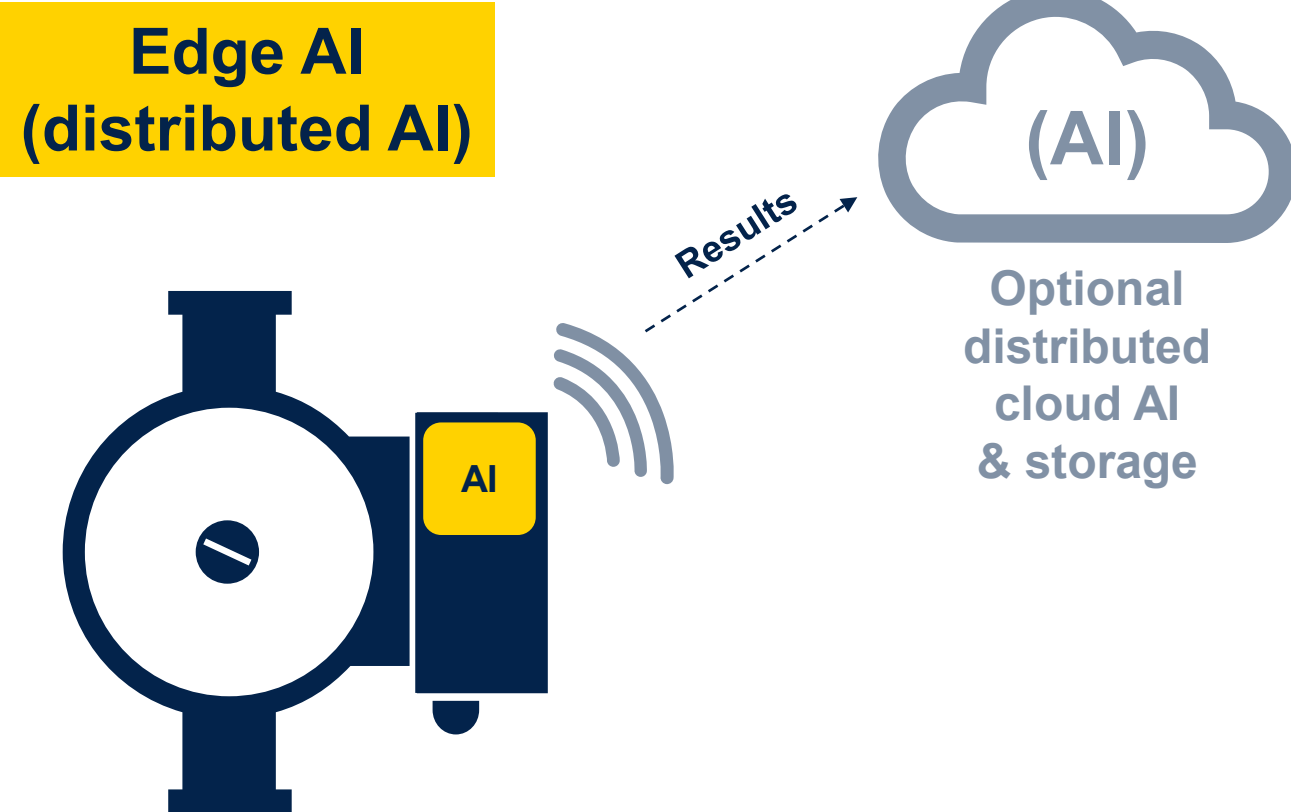
Global Data Generated Annually



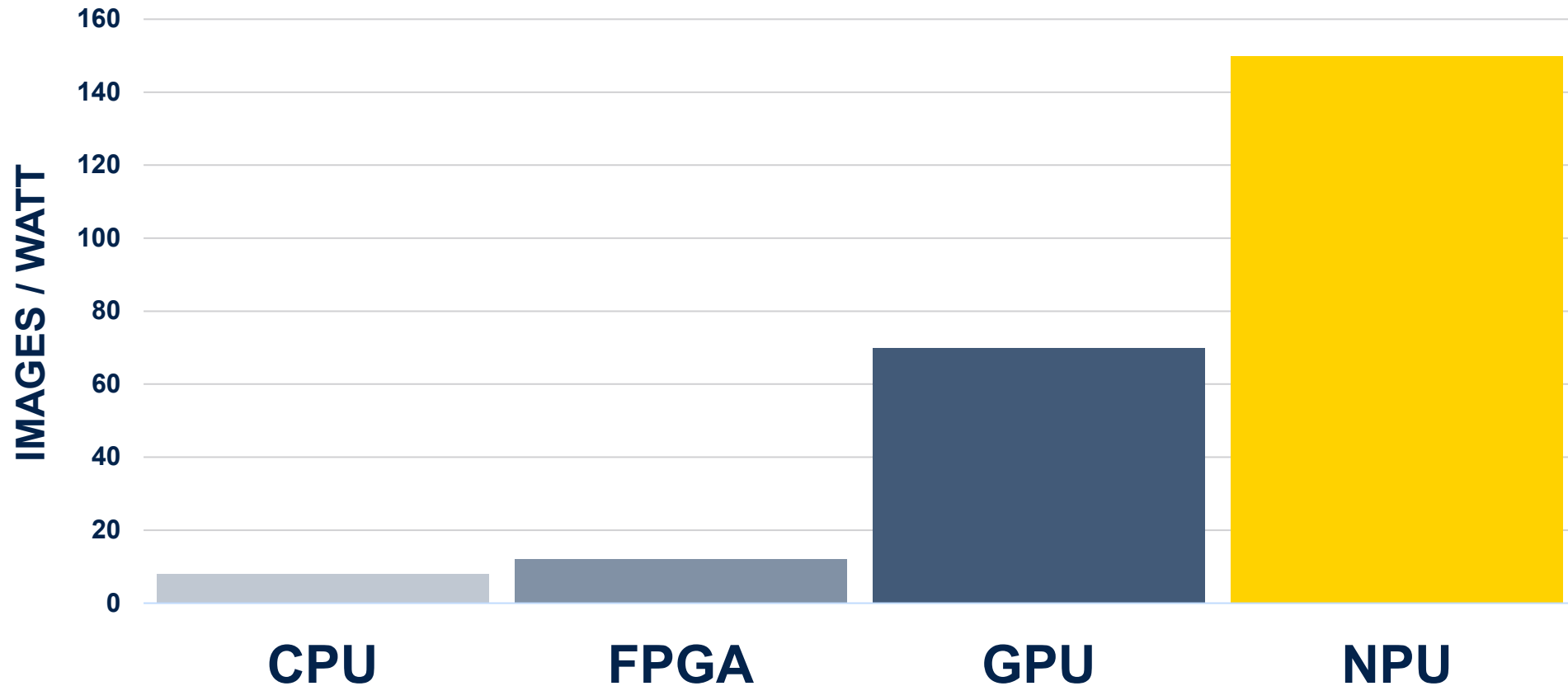
**120 ZetaBytes data generated in 2024
> 180 ZetaBytes in 2025**

Source: explodingtopics.com

The rise of edge AI: AI at device level



Edge AI acceleration requires new architectural solution: the NPU



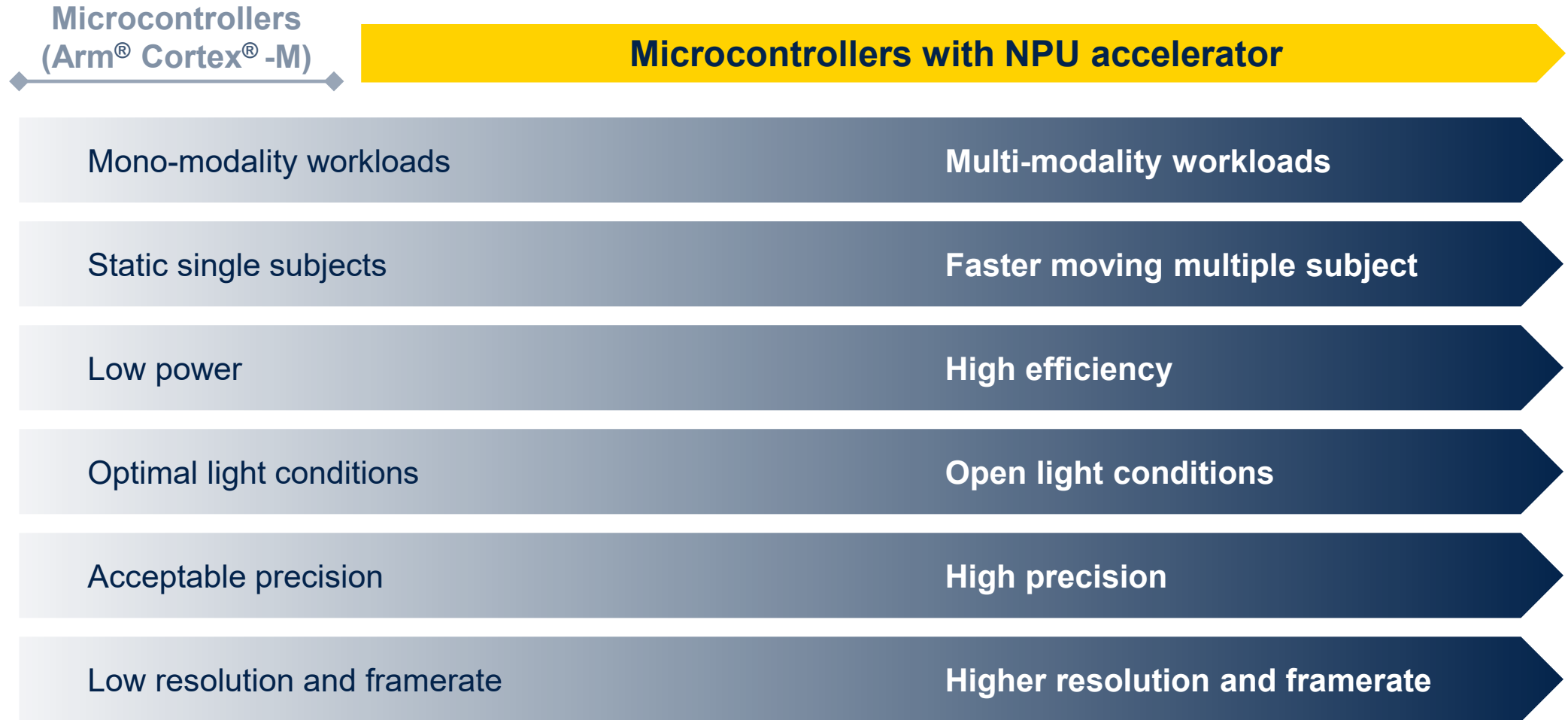
GPU: graphic accelerator

NPU: neural processing unit (AI accelerator)

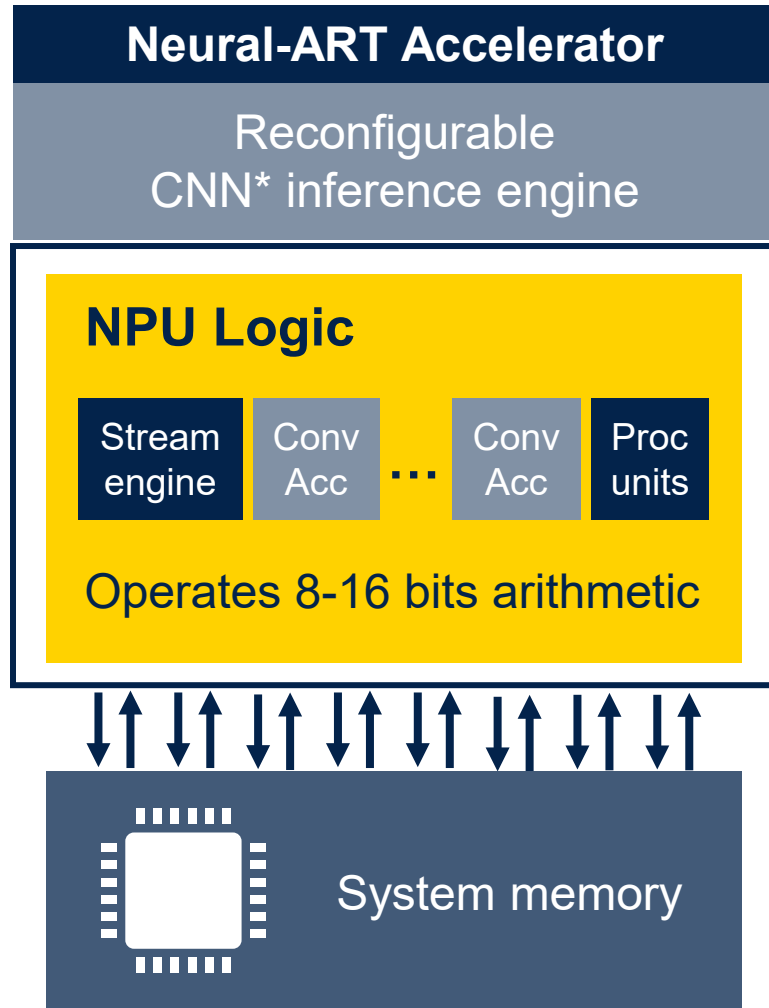
Source W. Dally

From DMIPS to TOPS, the paradigm shift Opening a new range of embedded AI applications

-  Object segmentation localization
-  Pose estimation
-  Object classification
-  Speech recognition
-  Sound analysis
-  Face/people detection
-  Wake word
-  Time series classification
-  Anomaly detection



Neural-ART Accelerator architecture overview



- **A paradigm shift** from the Von Neumann architecture towards a flexible, dedicated dataflow **stream processing engine**.
- Hardware acceleration for a **wide range of neural network architectures**.
- **Embedded security** to protect assets.
- **Seamless integration** into the MCU backbone via **two 64-bit AXI** interfaces.
- **Configurable** from 72 MACs to 2304 MACs.
- Achieves **up to 4.6 TOPS** at **1 to 5 TOPS/W****

* Convolutional neural network

** May vary according to technology node

Neural-ART Accelerator in STM32N6 MCU

600x
ML performance uplift*

Dedicated embedded neural processing unit

- 600 GOPS
- 3 TOPS/W power consumption
- Cache memory to optimize external memory access

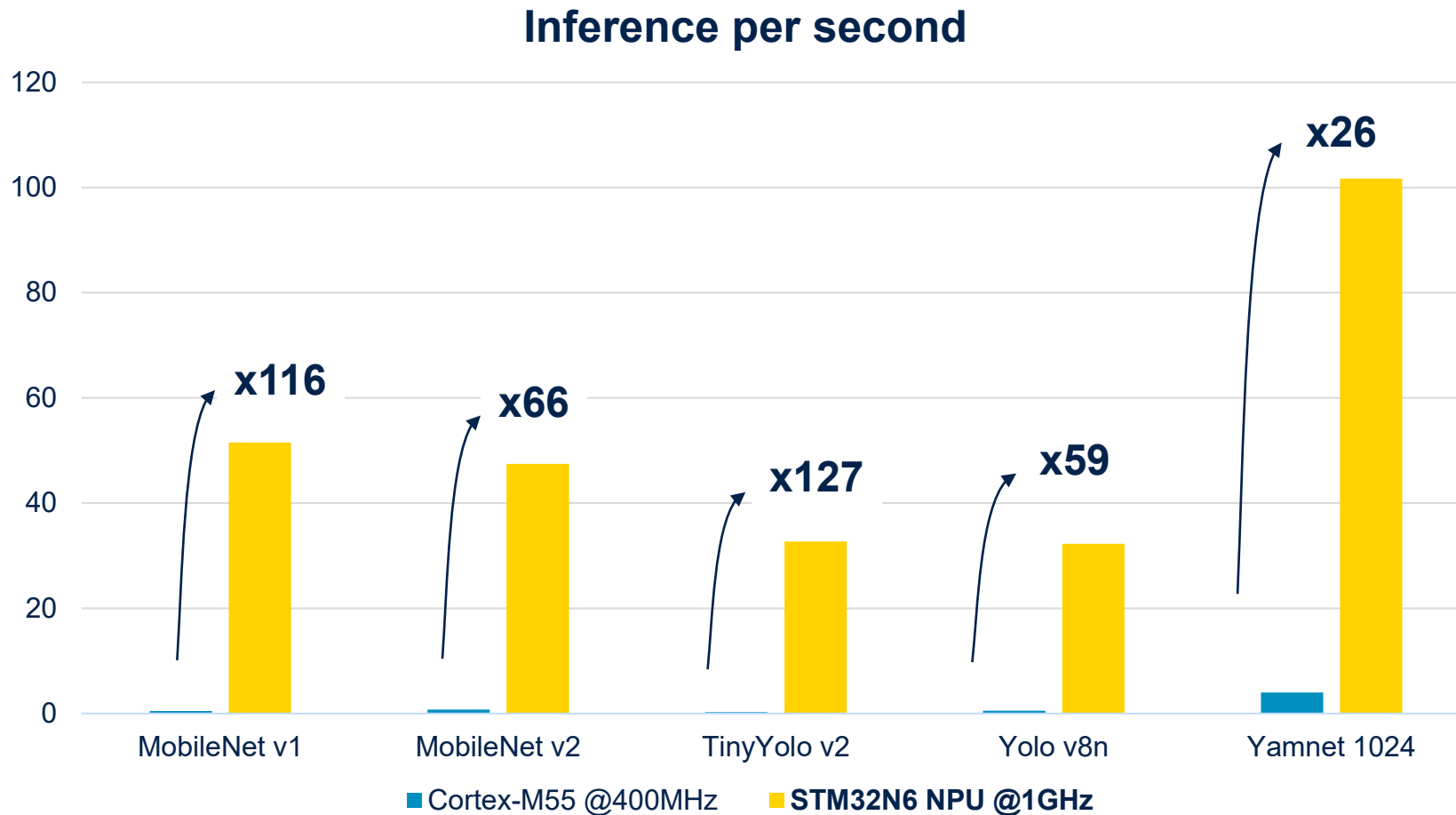
Dataflow stream processing engine
reduces MCU memory throughput
requirements and power consumption

Neural
processing
unit

MCU core

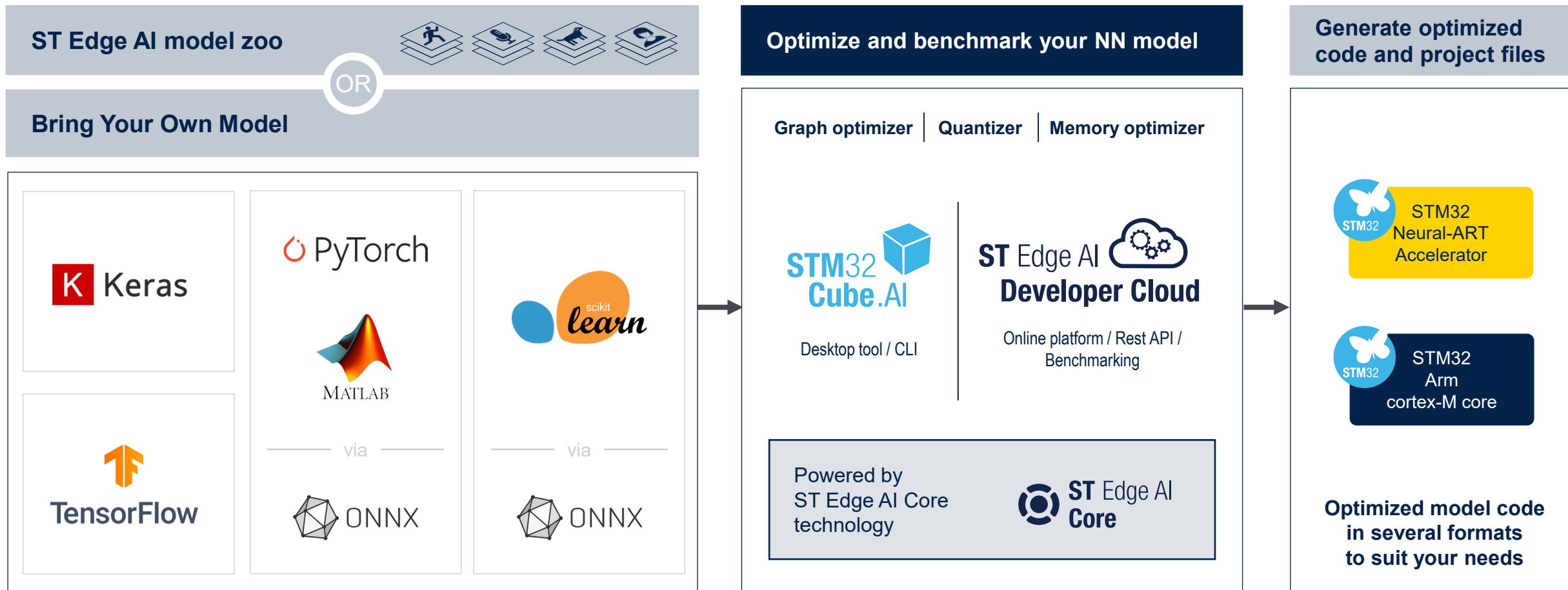
System
Memory

Neural-ART Accelerator provides a huge performance leap for AI inference



- **MobileNet v1:** image classification
- **MobileNet v2:** image classification
- **TinyYolo v2:** object detection
- **Yolov 8n :** object detection
- **Yamnet 1024:** audio recognition

Seamless integration with existing software ecosystem



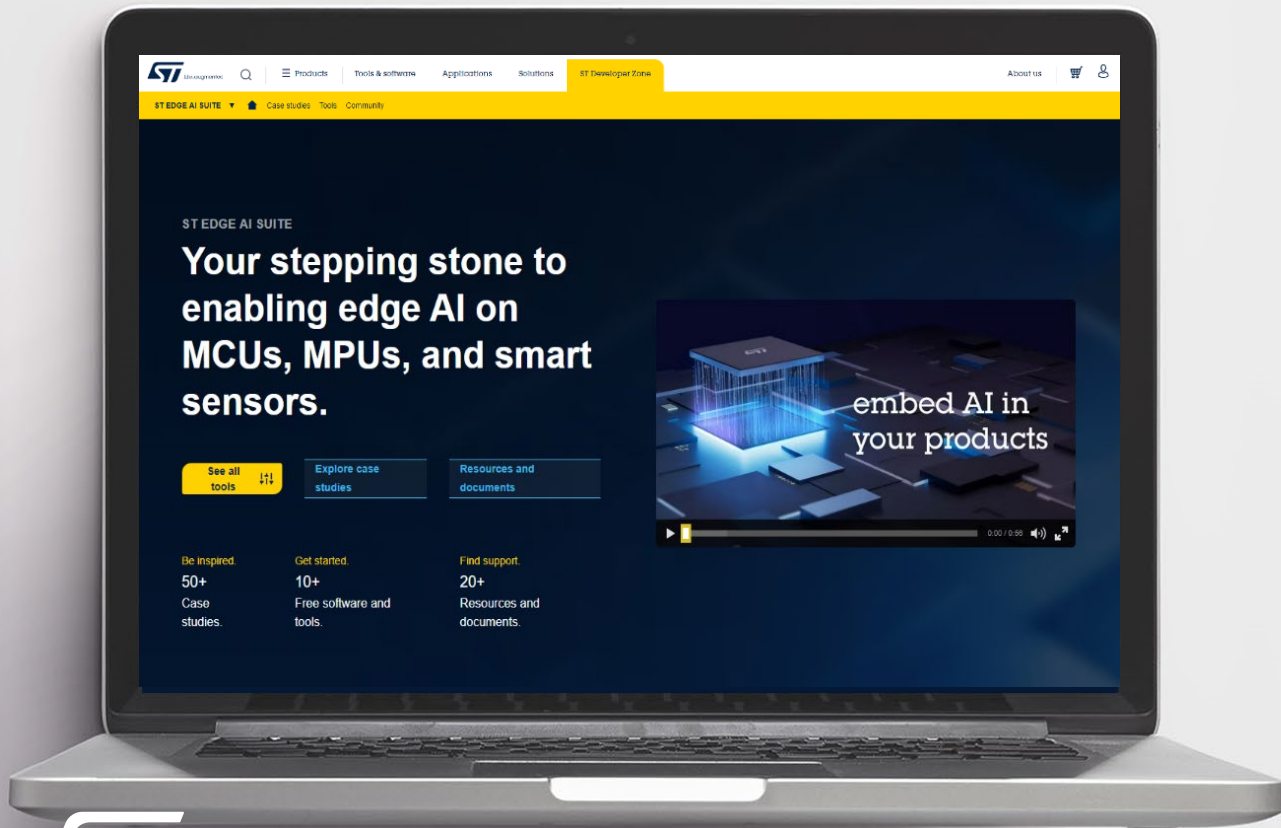
Reach the full potential of your application



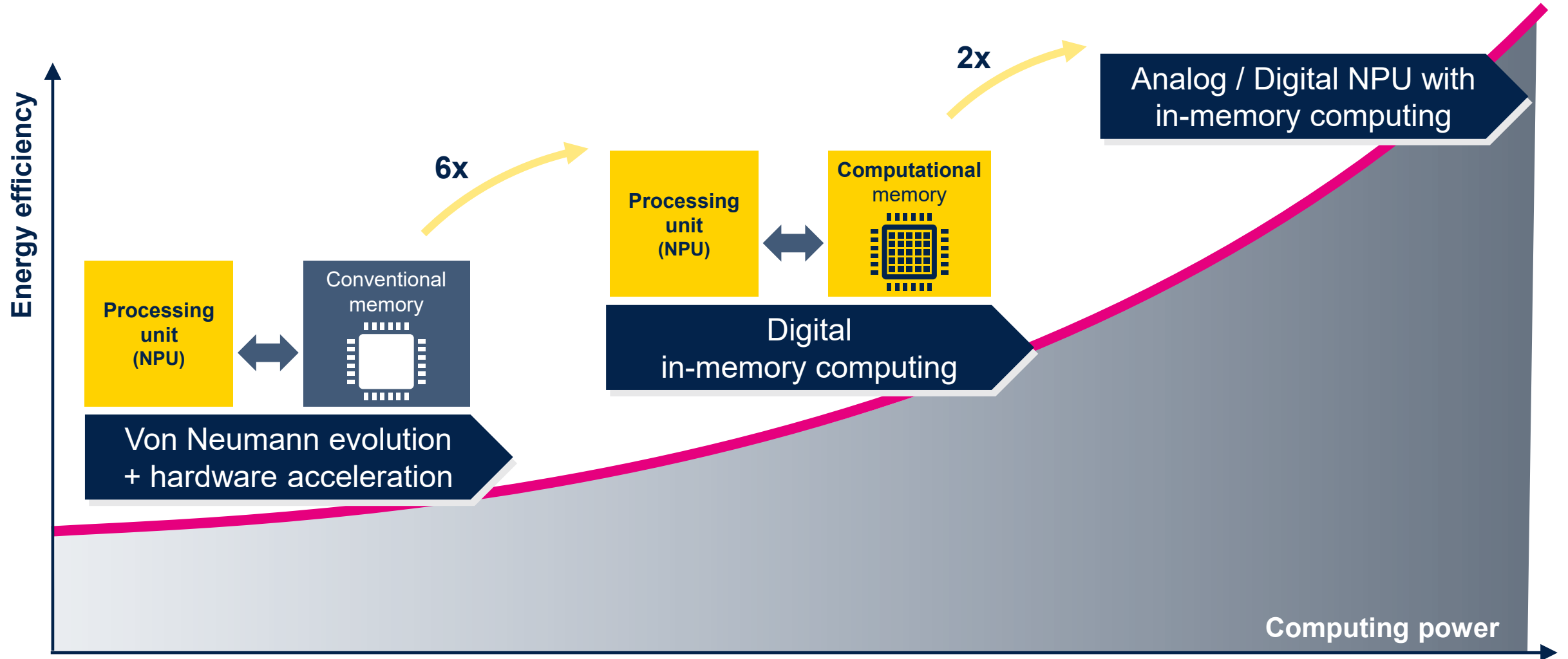
50+ case studies

10+ free software tools

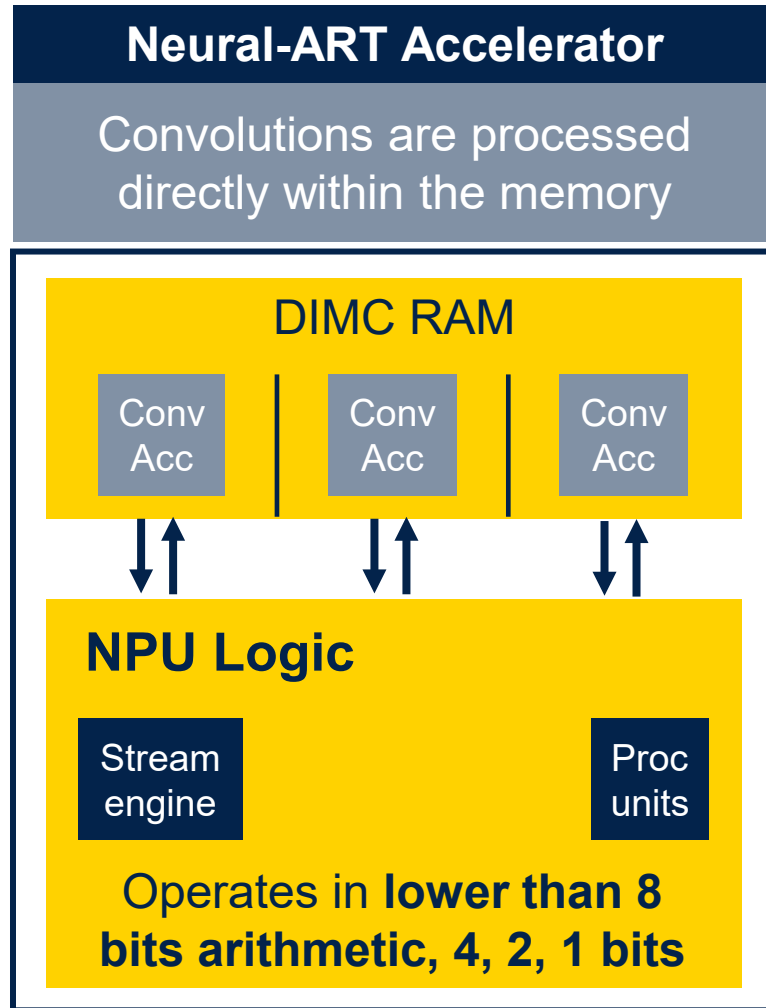
Unified ST Edge AI core technology



Neural-ART Accelerator outlook



More about next generations Neural-ART Accelerator



- **In-memory cell arithmetic** significantly reduces data transfer with memory hence power consumption.
- Achieves up to **6x improvement in TOPS and TOPS/W.**
- Support **advanced quantization (4, 2, 1 bit)** for further performance improvements.
- Ensures **seamless workflow integration** in the continuity of Gen 1.

Our technology starts with You



Read the whitepaper to know more

© STMicroelectronics - All rights reserved.

ST logo is a trademark or a registered trademark of STMicroelectronics International NV or its affiliates in the EU and/or other countries.

For additional information about ST trademarks, please refer to www.st.com/trademarks.

All other product or service names are the property of their respective owners.



life.augmented