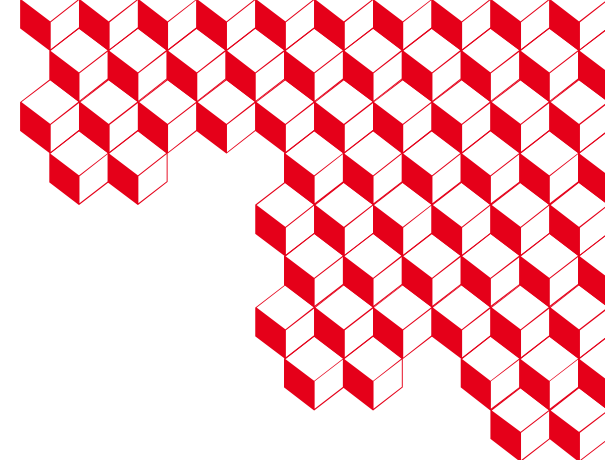




Les enjeux de l'IA embarquée

Olivier Bichler

olivier.bichler@cea.fr



Sommaire

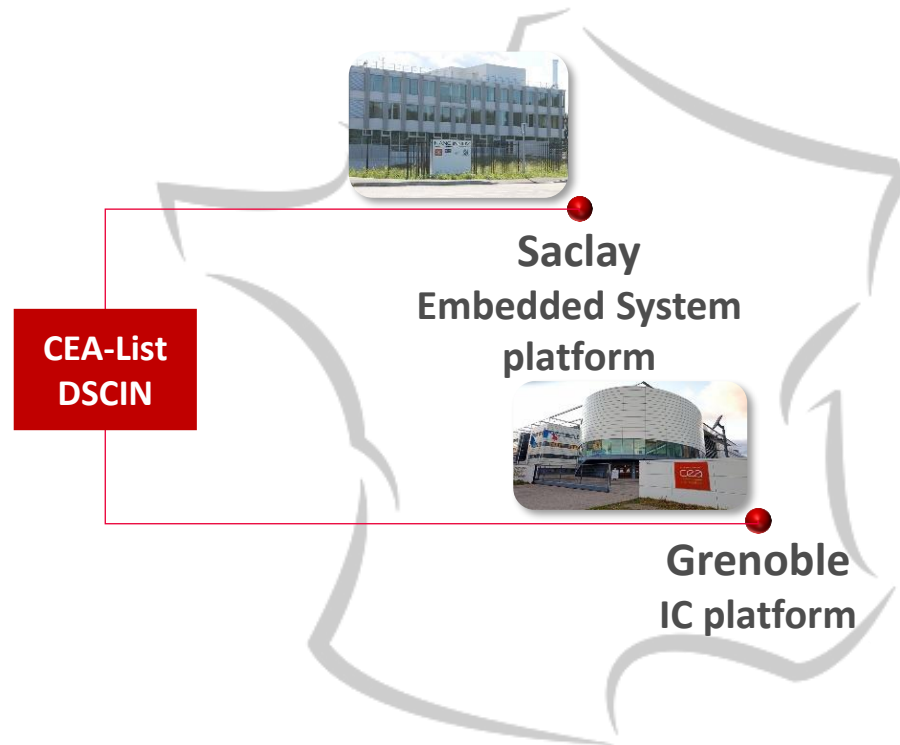
1. Introduction : le DSCIN et le LIAE

2. Les enjeux de l'IA embarquée

3. Nos travaux actuels



DSCIN: Digital System & Integrated Circuit division



200 Staff members

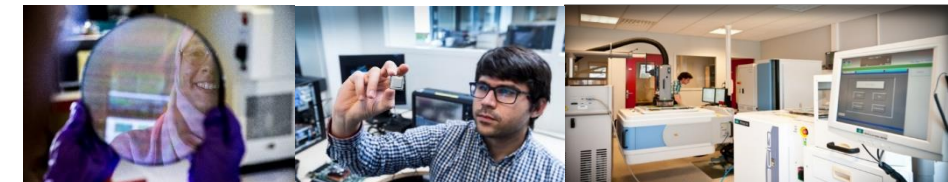
2 500 m² of facilities

30 Industrial partners

10 IC designs per year

25 Patents granted per year

100 Publications per year

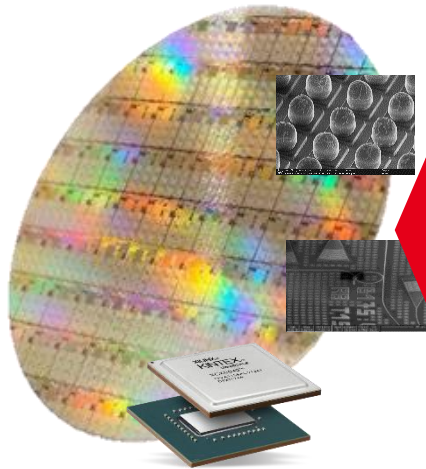


CEA Digital system design

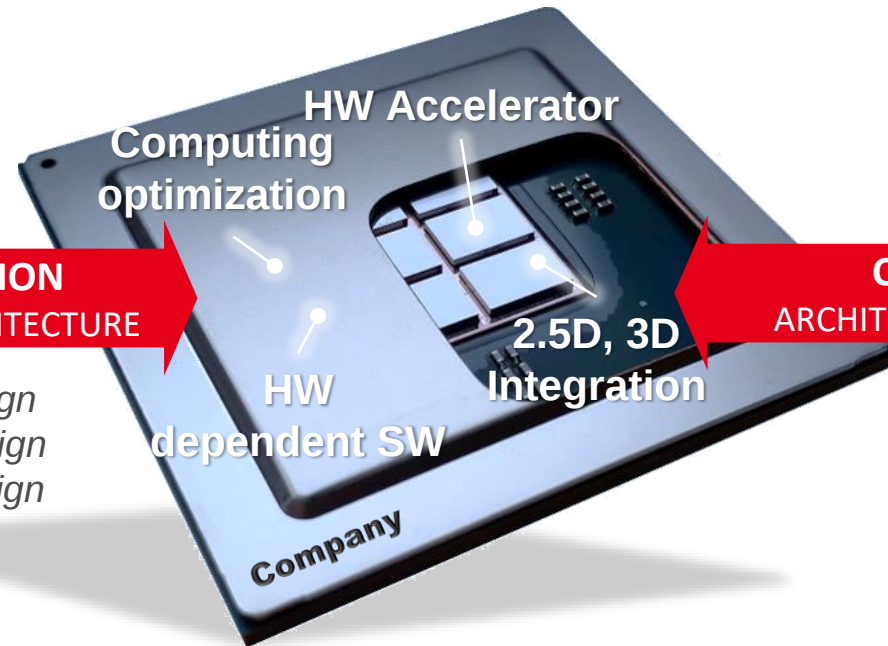


Eco-innovation

Semiconductor technologies



A 360° APPLICATION DESIGN



CO-OPTIMIZATION
TECHNOLOGY & ARCHITECTURE

- Automated design
- Full custom design
- Accelerator design

CO-DESIGN
ARCHITECTURE & SYSTEM

- Algorithm
- Software stack
- Data coding

Smart computing
systems

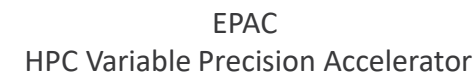
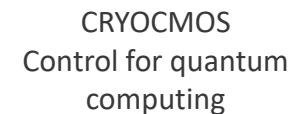
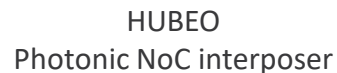
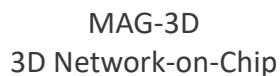


CHARACTERIZATION

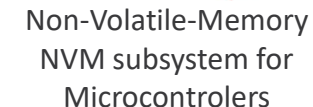
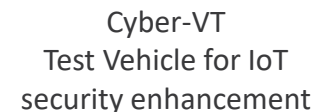
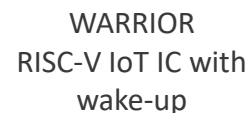
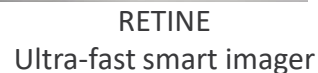
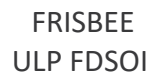
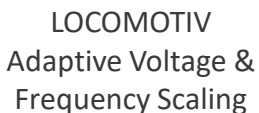
DESIGN

EXPLORE

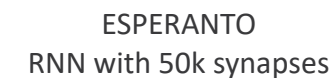
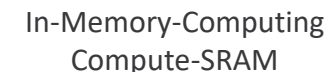
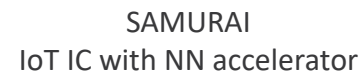
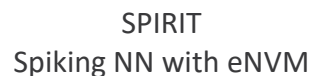
MODEL & SIMULATE

HI
EP

**THE
W**

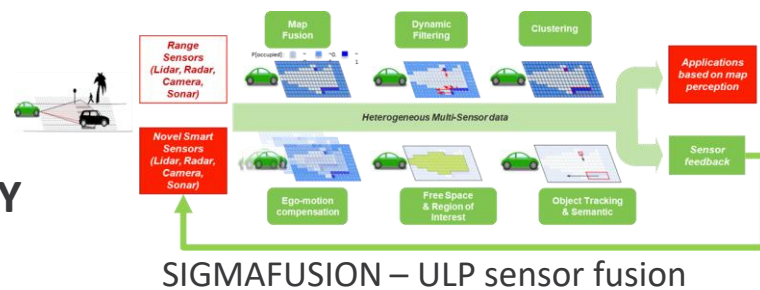


A

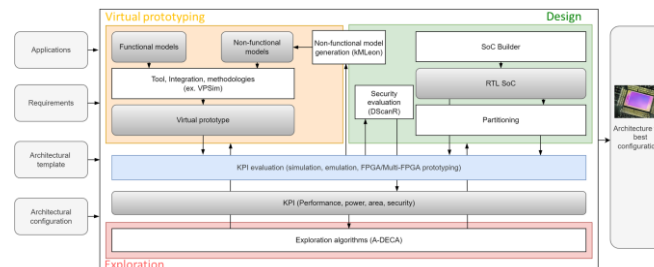


10+ years of experience in Systems co-design

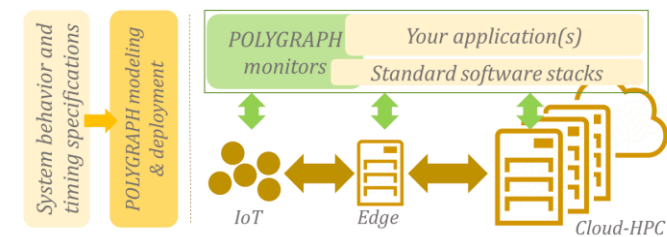
HIGH EFFICIENCY



SIGMAFUSION – ULP sensor fusion

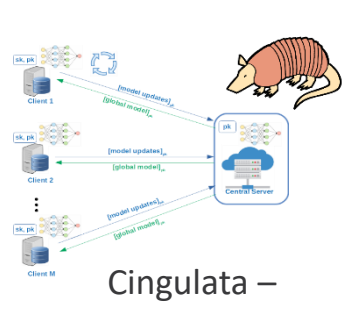


VPSIM – HW Digital Twin

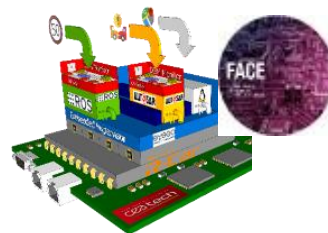


Polygraph- Edge2Cloud

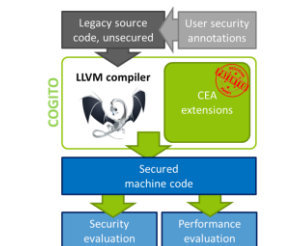
TRUST-WORTHY



Cingulata – cryptocomputing



MoCC – Communication & computing QoS



Cogito – SW Cyber

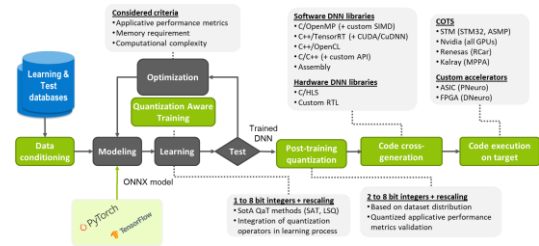


Leaf – real time reliability

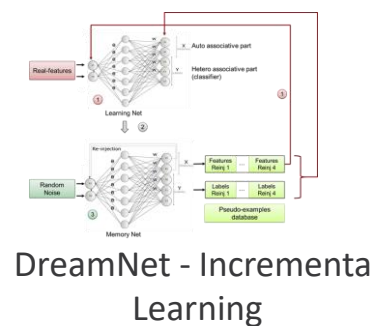
XANTHOS μKernel

AI

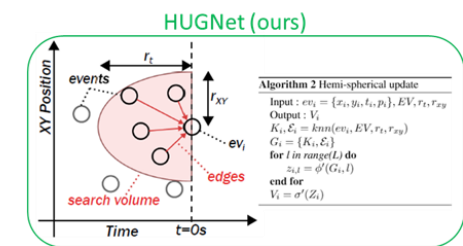
aidge



AIDGE embedded AI



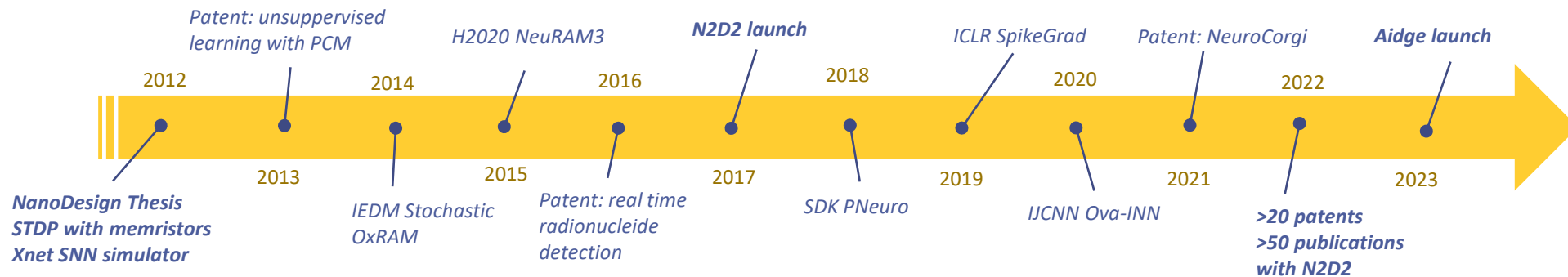
DreamNet - Incremental Learning



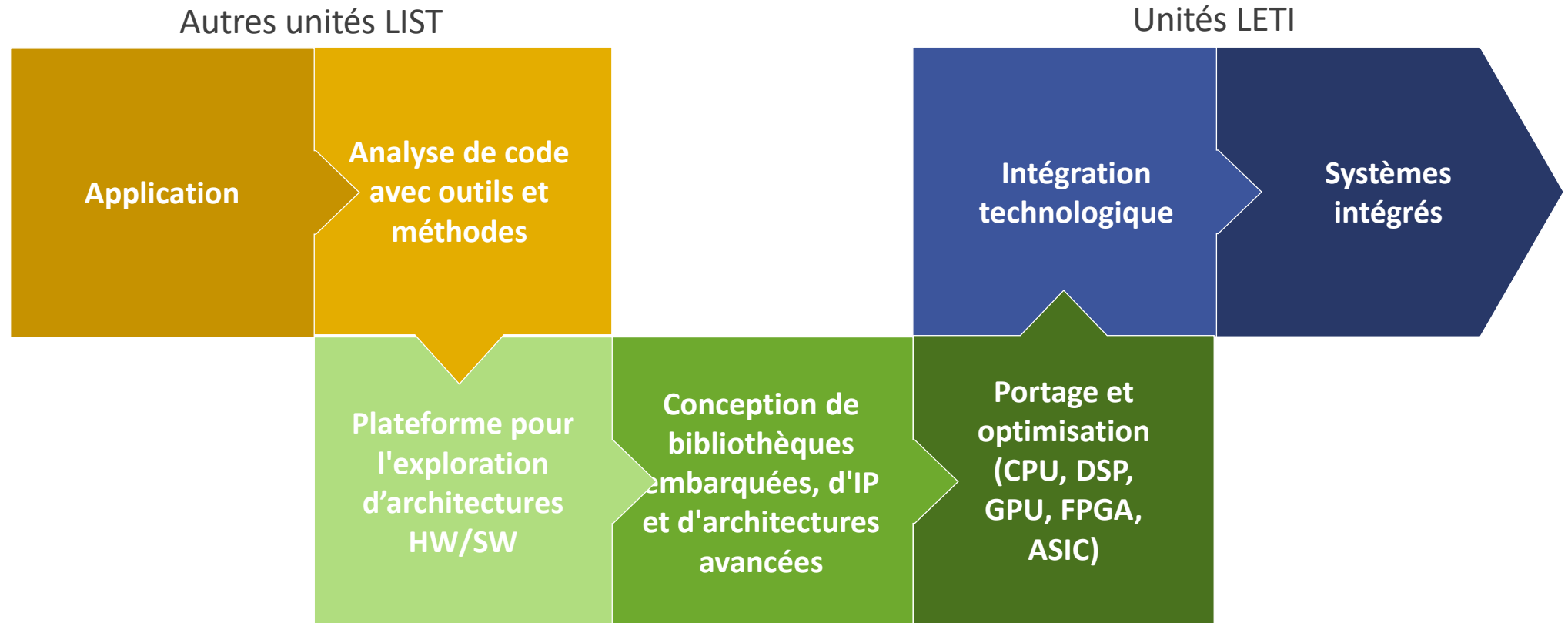
ULP GNN

Les origines du LIAE

Laboratoire Calcul Embarqué et Image <i>Jusqu'en 2009</i>	Laboratoire Calcul Embarqué <i>2010 - 2014</i>	Laboratoire Adéquation Algorithme Architecture <i>2015 - 2019</i>	Laboratoire Intelligence Artificielle Embarquée <i>Depuis 2020</i>
DTSI/SARC/LCEI	DACLE/LCE	DACLE/SCSN/L3A	DSCIN/LIAE
Mission : Concevoir des solutions HW/SW pour des systèmes embarqués complexes basés sur l'image. Les solutions sont principalement "classiques" avec un peu de recherche sur les réseaux neuronaux profonds.	Mission : Concevoir des solutions HW/SW pour des systèmes embarqués complexes. Séparation de l'équipe de traitement d'images de haut niveau avec la création du DIASI/LVIC.	Mission : Concevoir des solutions HW/SW pour des systèmes embarqués spécifiques à une application complexe. Séparation de la conception de l'architecture HW "générique" qui a créé le nouveau LCE.	Mission : Concevoir des solutions HW/SW pour des systèmes embarqués d'IA complexes. Intégration de l'activité des réseaux neuronaux profonds à la mi-2019.



LIAE est un joueur d'équipe clé



LIAE : Comblant le fossé entre le développement d'applications et l'intégration technologique

Sommaire

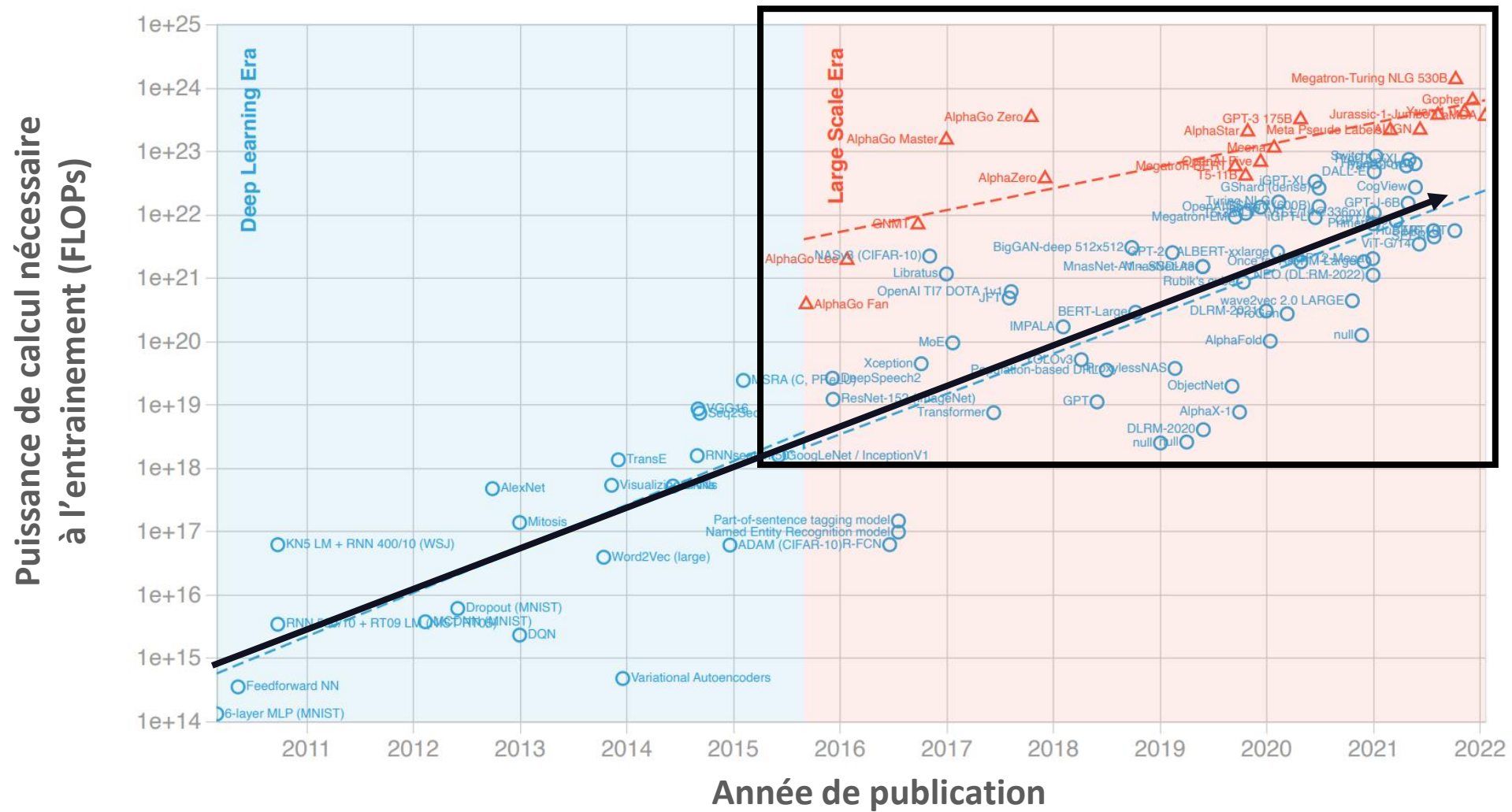
1. Introduction : le DSCIN et le LIAE

2. Les enjeux de l'IA embarquée

3. Nos travaux actuels



L'essor du *deep learning* et de l'IA



L'IA est largement dominé par le *deep learning*

L'essor de l'IA

T Qu'est-ce que chatGPT ? Réponds-moi en deux phrases.



GPT-3 175B Jurassic-1-Jumbo
AlphaStar Meta Pseudo Labels

Meena
Megatron-EMU
T5-11B
GPT-XL
GShard (dense)
OpenAI GPT-3 (336px)
Megatron-EMU
GPT-1-6B
PaLM
GPT-3.5

Switch
DALL-E
CogView
GPT-3.5

BigGAN-deep 512x512
GPT-2
ALBERT-xxlarge
Once for All
MnasNet-AM
SSD-Lite
Rubik's cube
NEO (DL-RM-2022)

AlphaZero
AlphaGo Lee
AlphaGo Fan
NAS-3 (CIFAR-10)
Libratus
OpenAI TI7 DOTA 1v1
JFT
MoE
Xception
MSRA (C, P, F, D)
DeepSpeech2
ResNet-152 (ImageNet)
Transformer
GoogLeNet / InceptionV1

AlphaFold
wave2vec 2.0
LARGE
DLMR-2020
null

ObjectN

DLRM-2020
null

Alpha

DLRM-2020
null

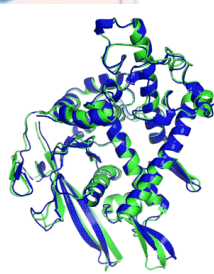
DLRM-2020
null

DLRM-2020
null

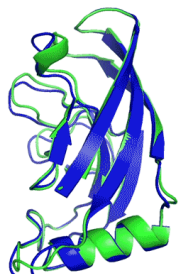
DLRM-2020
null

DLRM-2020
null

DLRM-2020
null

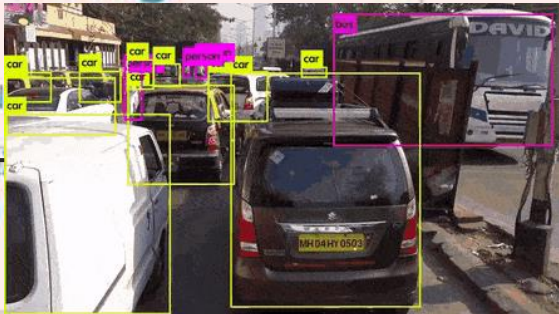
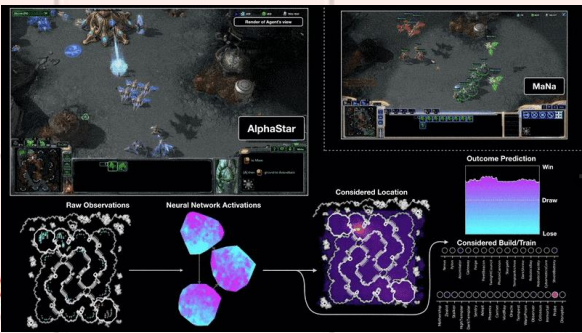


T1037 / 6vr4
90.7 GDT
(RNA polymerase domain)



T1049 / 6y4f
93.3 GDT
(adhesin tip)

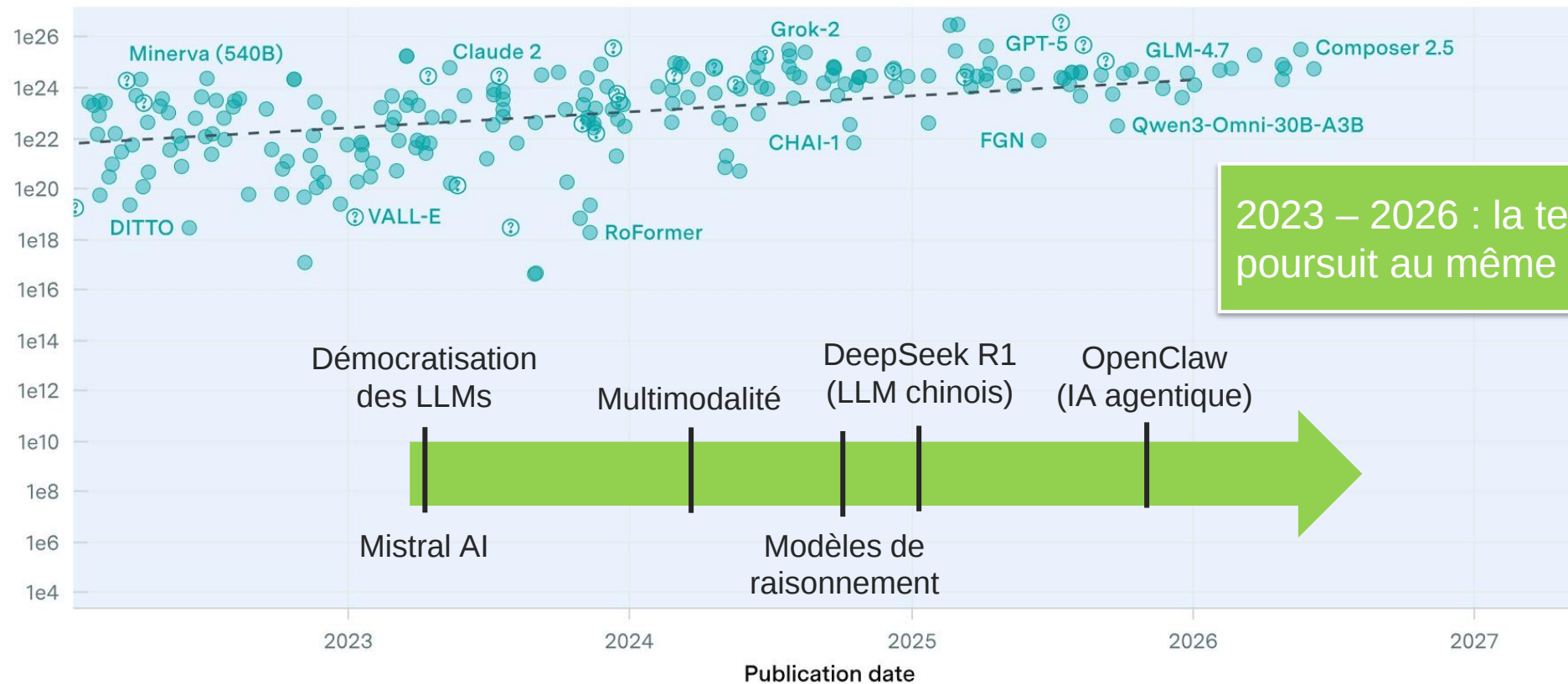
● Experimental result
● Computational prediction



L'essor de l'IA

Training compute (FLOP)

?: Speculative data 552 Results



Un développement aujourd'hui exponentiel

Bank of America: "It took ChatGPT just 5 days to reach 1 million users, 1 billion cumulative visits in 3 months and an adoption rate which is 3x TikTok's and 10x Instagram's,"

"The technology is developing exponentially."

Marché IA adressable selon
SpaceX d'ici 2030 : **26 000Md**
CA visé en 2030 : 322Md

Croissance globale du marché de l'IA : de 200Md en 2023 à 2 000Md¹

■ de l'edge IA : de 10Md à 100Md² en 2030 → x10

Explosion des investissements

- Microsoft / OpenAI : 10Md en 2023, supercalculateurs à 10k-100k GPUs, entre 100k-1M GPUs déployés⁴
- Google : des douzaines de supercalculateurs à 4k TPUs v4⁵
- Twitter : 10k GPUs
- France/GENCI : Jean Zay (3k GPUs), CEA-List/Factory-IA (1k GPUs)...

xAI Colossus : 200k H100 GPUs

Mistral :

- 2026 : mise en service d'un centre de données à Bruyères-le-Châtel
- 1 GW de capacité en 2030

L'IA pourrait contribuer à 20% de croissance du PIB de l'Europe en 2030, soit 3 000Md³

Soutenu par des plateformes *open source*

Des outils performants et accessibles



Les US dominent le marché avec des plateformes de deep learning largement utilisées dans le monde positionnées dans le cloud aujourd'hui



GAFAM



La Chine entre dans la course avec ses propres plateformes de deep learning issues des BATX. Les autorités chinoises investissent aussi massivement sur le sujet



Plan AI+ lancé en 2024

- Modèles de fondation
- Infrastructure souveraine
- Processeurs IA souverains

15 000 robots humanoïdes produits en 2025 (90% de la production mondiale)

INVESTISSEMENT PUBLIC :

- National Science and Technology Council
- DARPA Defense Advanced Research Projects Agency
- Advanced Research Projects Agency-Energy

INVESTISSEMENT PUBLIC :

- Les autorités chinoises ont annoncé qu'elles investiraient jusqu'à 59 milliards de \$ d'ici à 2025 dans l'objectif de rattraper les États-Unis

Solutions profondément orientés **déploiement GPU/Cloud** (quasi-monopole NVidia)

Jusqu'où ?

L'IA pourrait représenter plus de la moitié du *cloud computing* en 2030

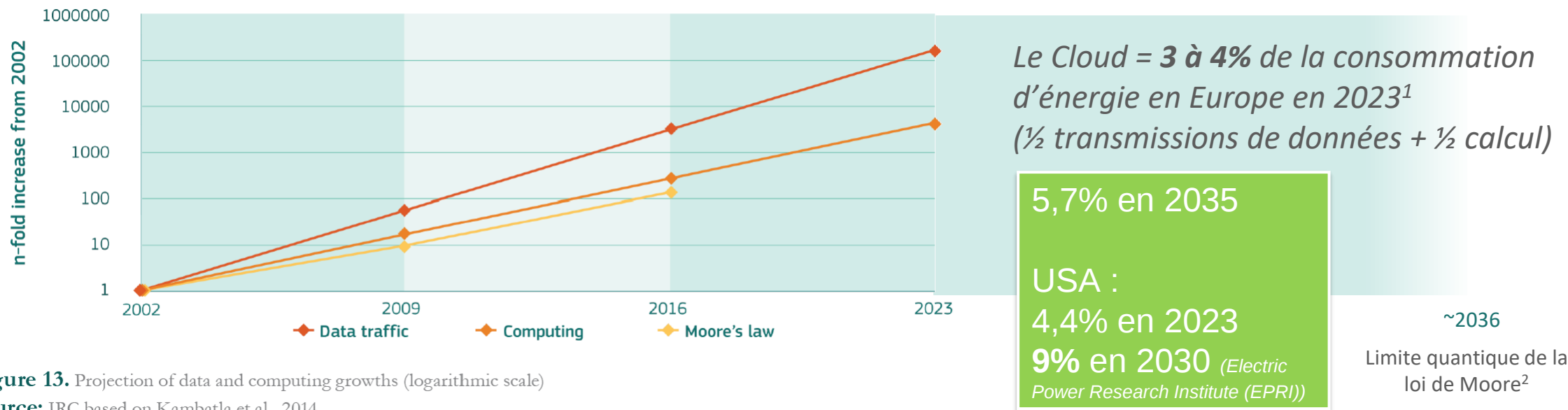


Figure 13. Projection of data and computing growths (logarithmic scale)

Source: JRC based on Kambatla et al., 2014

*The expected growth in data, computing processing, and energy consumption are changing the present paradigm "do everything in the cloud" to a new one that can be called "**bring computing to where the data are**".*

L'edge IA, une nécessité pour **passer à l'échelle et diminuer le coût !**

Vers l'essor de l'IA embarquée

Le déploiement, du Cloud vers l'*edge* :

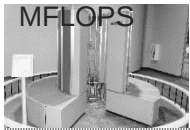
- Maturité des applications
- Passage à l'échelle
- Diminution du coût
- Interaction avec le monde physique

- Souveraineté
- Confidentialité

IBM ASCI White
7.23 TFLOPS



Cray-1
160
MFLOPS



ARM Cortex-M4
250 MOPS



Apple A16
17 TOPS



Taalas (start-up à Toronto)
Llama sur silicium atteint
17 000 jetons/sec !
10x meilleur que le SotA

Limite quantique de la
loi de Moore²

Cloud IA non soutenable

Cloud IA : développement
(apprentissage large échelle,
génération de circuits)

Ere des circuits « câblés »

Edge IA :
déploiement
(inférence & adaptation)



Smart Cities



Industrial Robotics



Autonomous Vehicles



Warehouse Robots



Drones



Robotic Cleaners



Augmented Reality



Smart Agriculture

Puissance de calcul (TOPS) – échelle log

1985

2000

2010

2022

~2036

Vers l'essor de l'IA embarquée



La tendance de fond du développement du Edge computing, ou Intelligence Artificielle Embarquée, représente une valeur potentielle considérable estimée entre 175 à 215 milliards US\$ en 2025 pour la seule composante hardware.

Une prévision commode mais « trop » visionnaire ?

« L'IA embarquée promet donc (défense, agriculture...), et repré

Des catalyseurs forts :

- Capacités autonomes dans la défense (guerre en Ukraine, instabilités...)
- Robotique humanoïde industrielle et domestique (Plan AI+ Chine, 1X, Tesla...)
- Banalisation des NPUs compatibles LLM/VLM *low cost* (notamment chinois : Sophon, SpaceMIT, RockChip, Huawei...)

L'USINE NOUVELLE

« L'edge computing est une plus grosse opportunité que le cloud », Thierry Breton, PDG d'Atos (2019)

« Le client nous demande de les traiter dans les datacenters [...] on arrive à un volume de 40 zetaoctets, soit 40 000 milliards de milliards de données. Selon différents cabinets, 80% vont être traités de façon centralisée sur les datacenters et le cloud, et 20% ailleurs, à la périphérie ou l'edge [...]. A l'horizon de cinq ans, ce seront 80% à traiter dans l'edge et 20% dans les datacenters et le cloud. »

Pourquoi embarquer l'IA ?

Motivateurs :

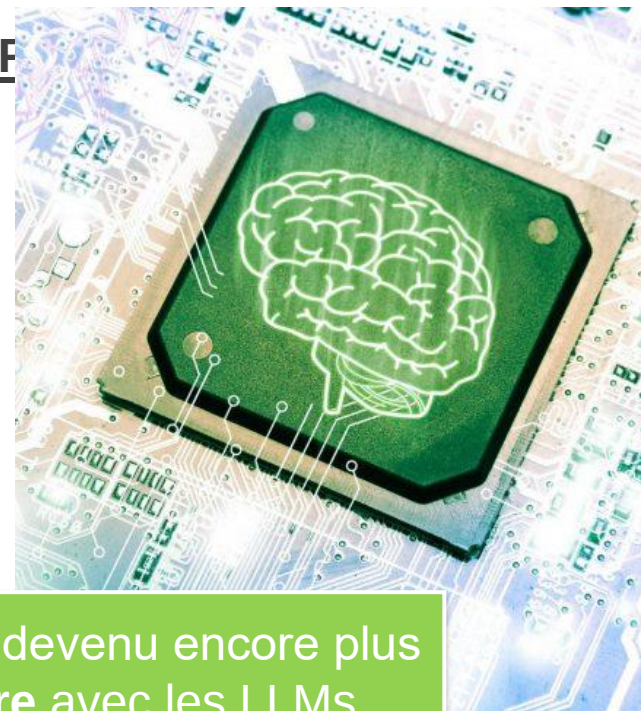
Passage à l'échelle & réduction des coûts
(transmission, infrastructure...)

Temps réel « dur »
(temps de réponse garanti)

Sûreté de fonctionnement
(disponibilité garantie)

Confidentialité des données
(si on est obligé...)

Oui, c'est devenu encore plus
nécessaire avec les LLMs,
et ce à tous les niveaux
(états, industries, individus)



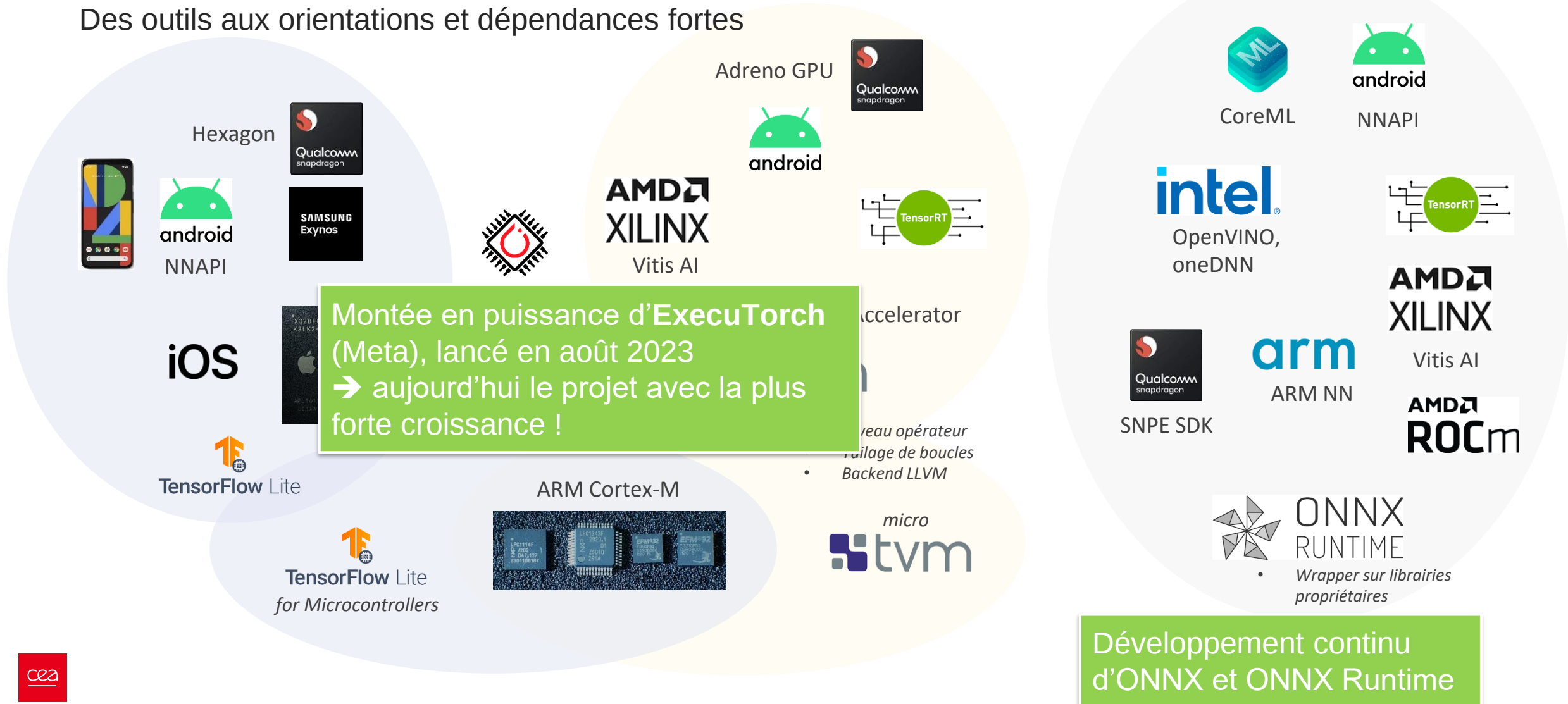
Fiabilité : capacité à garantir
les résultats
(explicabilité, robustesse...)

Frugalité : capacité d'adaptation
avec peu de données
(*zero-/one-shot learning*,
continual learning...)

Importance de la
mémorisation du contexte
pour les grands modèles

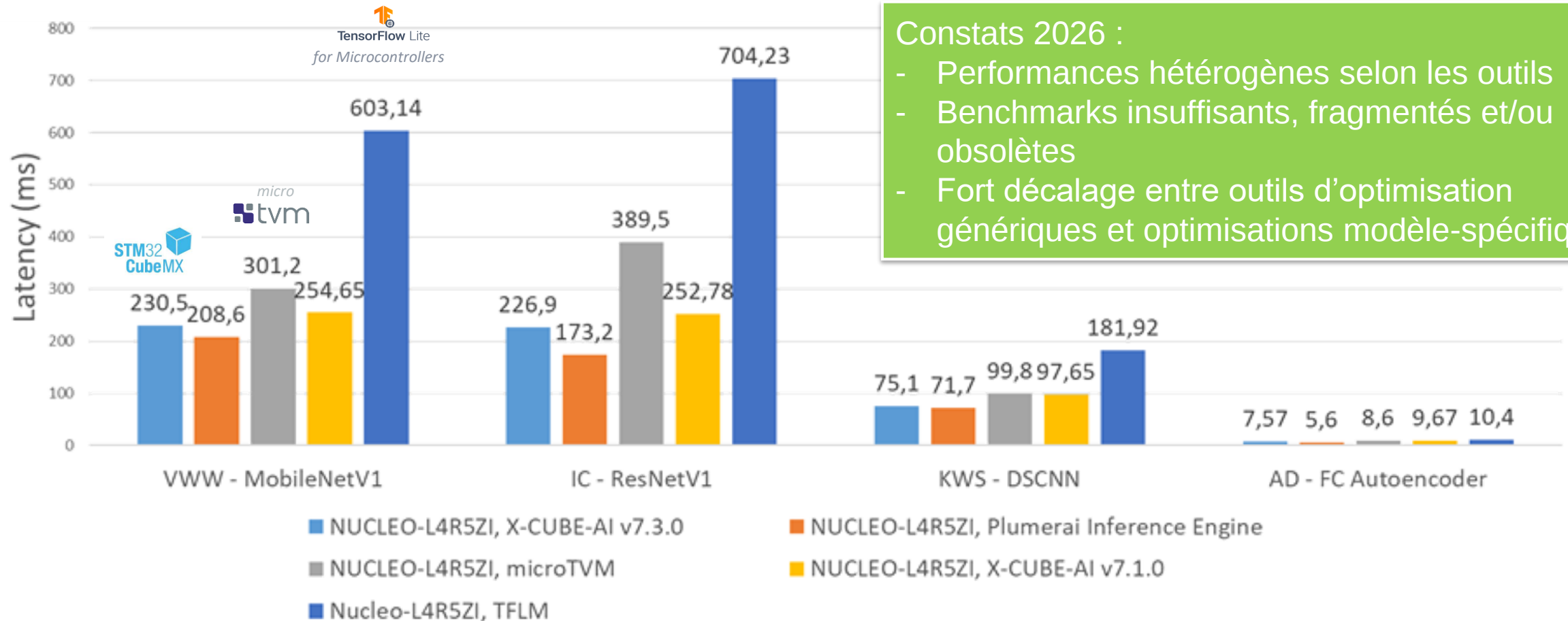
Quelles solutions pour l'IA embarquée ?

Des outils aux orientations et dépendances fortes



Quelles solutions pour l'IA embarquée ?

Pas forcément optimisés : benchmark 2023 MLPerf Tiny sur Cortex-M4

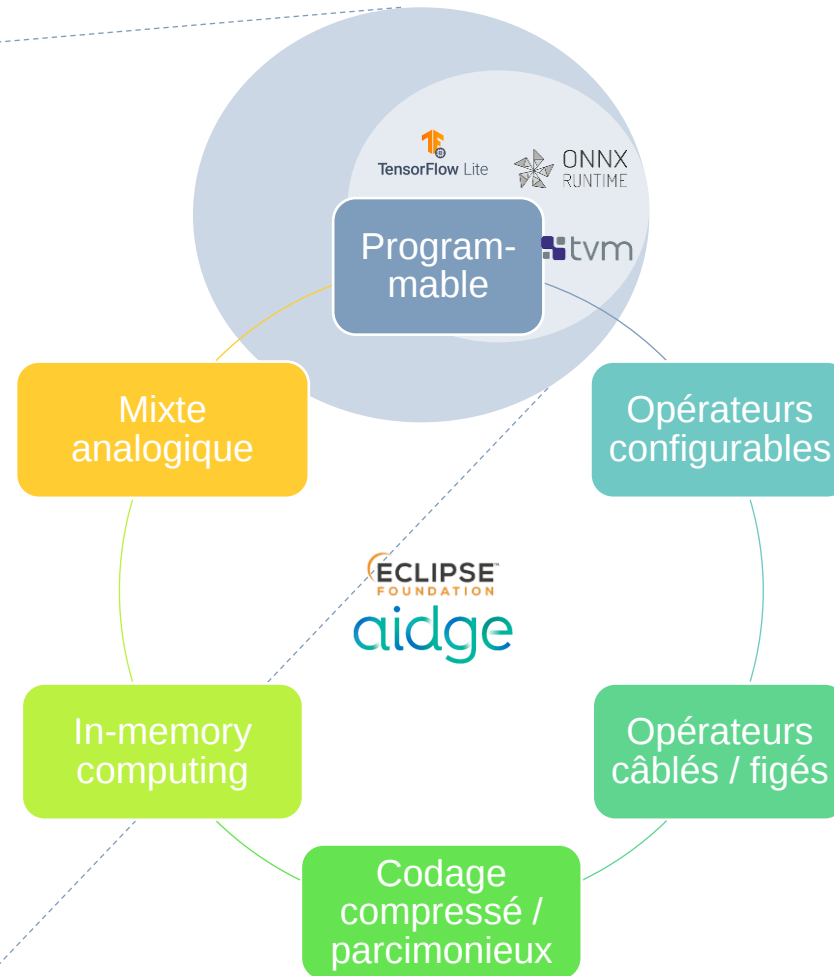


Quelles solutions pour l'IA embarquée ?

Contrastant avec la diversité des solutions et architectures matérielles

Processor	Company	Year	Process
ERGO	Perceive	4Q20	22nm
Recogni	Recogni Inc.	4Q22	7 nm
ARC NPX6-96K	Synopsys		
AiWare 4	AiMotive		
AiWare 3	AiMotive		
NeuPro-M	CEVA		
NeuPro-S NPS4000	CEVA		
Tyr 3	Vsora		
Tyr 1	Vsora		
AD1028	Vsora		
Origin E4	Expedera		
RAPID	IBM		
TruStream	Cornami		
NNE 110	Cadence		
NEOX A	Think Silicon		
SAKURA DNA-A800	EdgeCortex		
Journey 6	Horizon Robotics		

Processor	Company	Year	Process
Journey 5	Horizon Robotics	4Q21	
Ethos-N78	ARM	2Q21	7 nm
Hailo-8	Hailo	2Q20	16 nm
NMP 700	AiM Future	1H23	
Tiny Raptor v2	Dolphin	1H22	22 nm
Chimera QB4	Quadric	1Q23	7 nm
Q16	Quadric	3Q21	16 nm
Myriad X	Intel Movidius	4Q18	16 TSMC
Infer X X1	Flex Logix	1Q23	16 nm TSMC
WARBOY	FuriosaAI	1Q23	14 nm
ARC EV74	Synopsys	1Q20	16 nm



Processor	Company	Year	Process
Spiking neural Processor			
Loihi 2	Intel	4Q21	Intel 7
In memory DLA			
AI-X Lightspeed 2803S	Gyr Falcon Technology	4Q20	28nm
MX3	MemryX	2H22	
SpeedAI240	Untether AI	3Q23	7 nm TSMC
RunAI200	Untether AI	2Q21	16 nm
Analog DLA			
AML100	Aspinity	4Q22	
M1076	Mythic	1Q22	40 nm
GPX-10	Ambient	2Q21	40nm
Passage	Lightmatter	4Q22	
Mars	Lightmatter	4Q21	90nm (photo die) 14nm (dig. die)
Aurora DeepLight 1.0	Cognifiber	2H23	

Les enjeux pour l'IA embarquée

Au-delà des **performances, indispensables**, de **multiples enjeux techniques** pour assurer notre compétitivité :

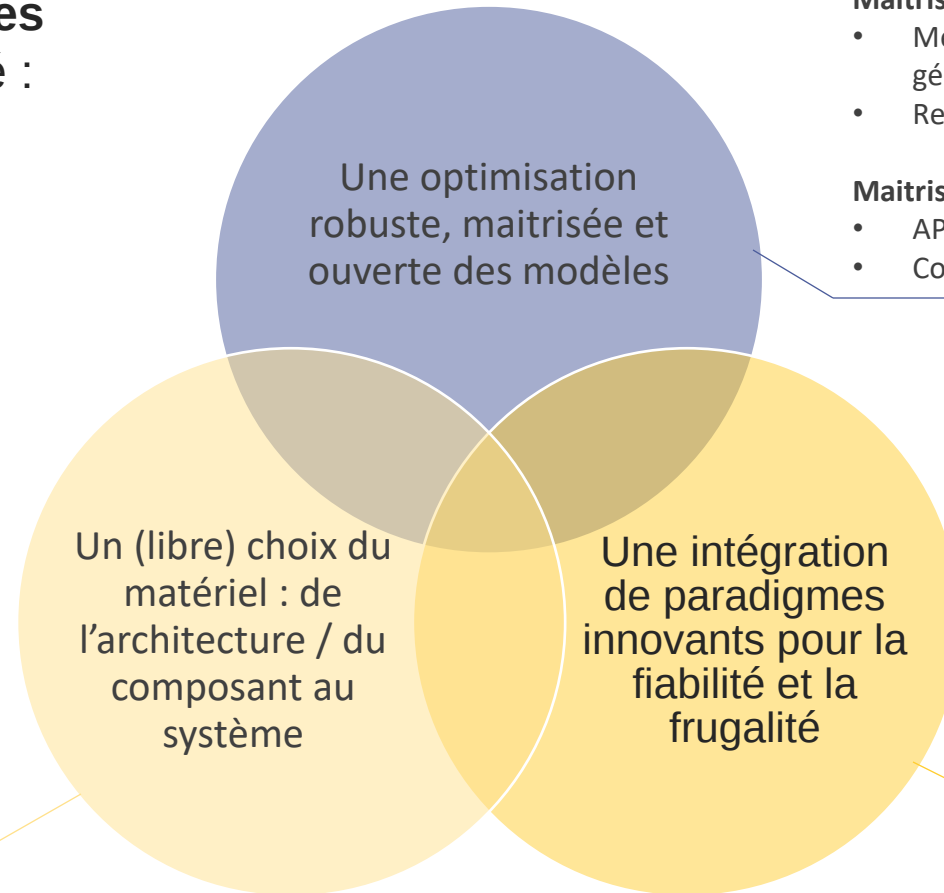
- *Maitrise*
- *Souveraineté*
- *Innovation*

Hétérogénéité :

- Multi-cibles matérielles
- Gestion de systèmes hétérogènes
- Multi-paradigmes de calcul

Souveraineté :

- Noyaux de calcul ouverts et indépendants
- Architectures et composants souverains
- Intégration de technologies souveraines (FDSOI, 3D-chiplet, mémoires NVM)



Maitrise des algorithmes :

- Méthodes ouvertes, documentés, génériques et réutilisables
- Reproductibilité des modèles optimisés

Maitrise du déploiement :

- API modulaire et unifiée pré- et post-déploiement
- Compatible avec des objectifs de certifiabilité

Fiabilité :

- Intégration de connaissances « innées »
- Maitrise du « domaine de vol »
- Robustesse aux fautes et aux attaques

Frugalité :

- Transférabilité de modèles « maitres »
- Adaptation hors ligne et en ligne

Sommaire

1. Introduction : le DSCIN et le LIAE

2. Les enjeux de l'IA embarquée

3. Nos travaux actuels

Modèles IA d'intérêt

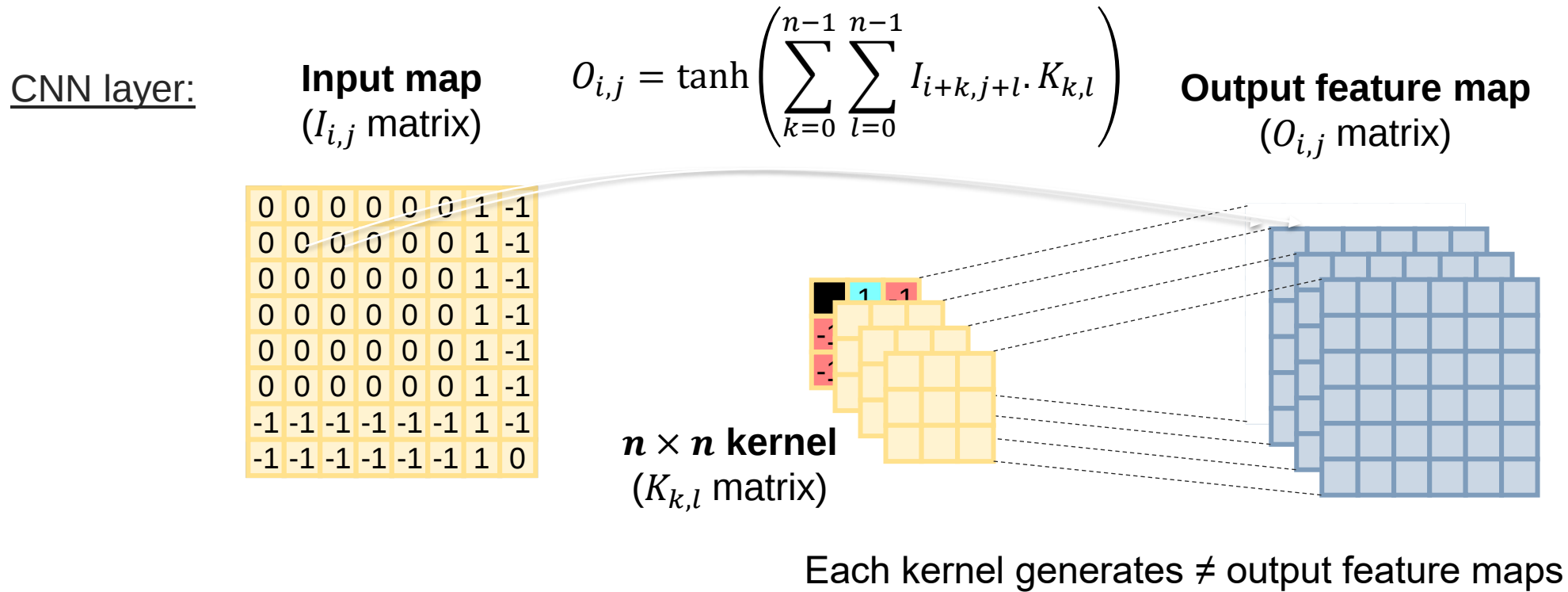
Plateforme logiciel open source pour l'IA embarquée : **Aidge**

Exemple de circuit ASIC ultraspécialisé : **NeuroCorgi**

Vision du laboratoire



Convolutional neural networks overview

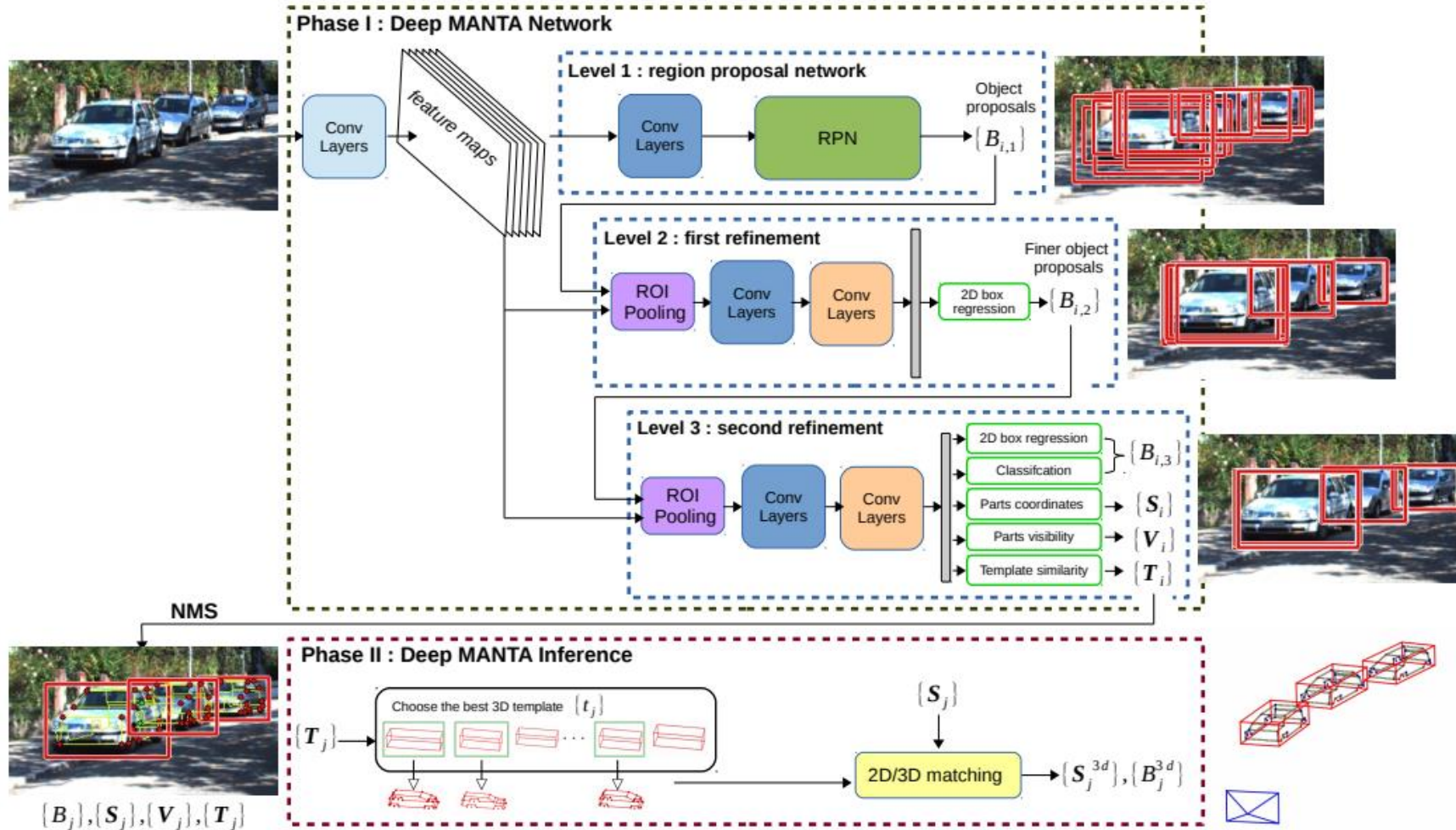


Convolution operation:
$$O_{i,j} = \tanh \left(\sum_{k=0}^{n-1} \sum_{l=0}^{n-1} I_{i+k,j+l} \cdot K_{k,l} \right)$$

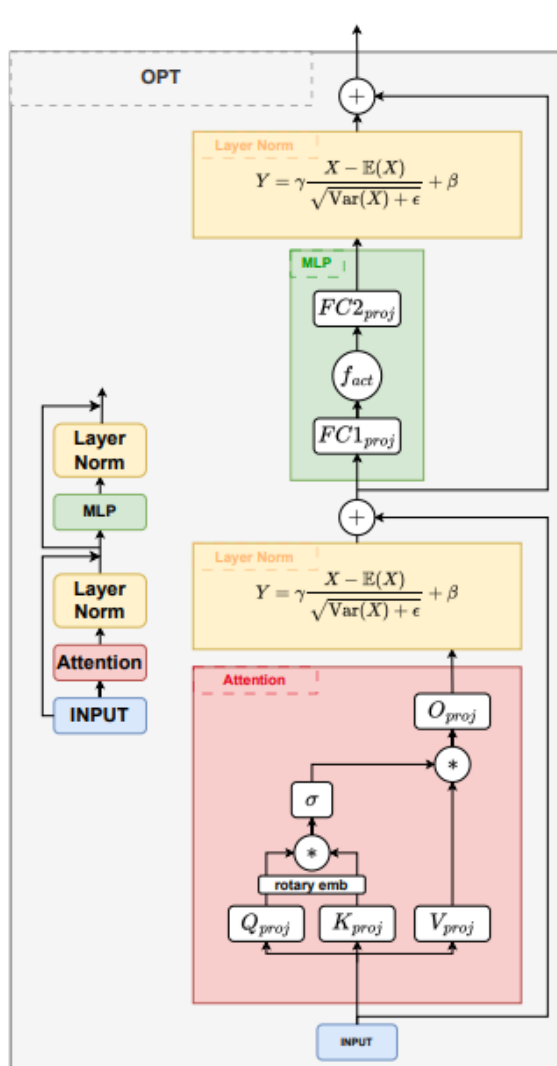
Kernels are learned with gradient-descent algorithms
(classical back-propagation is very efficient!)

Complex neural network systems

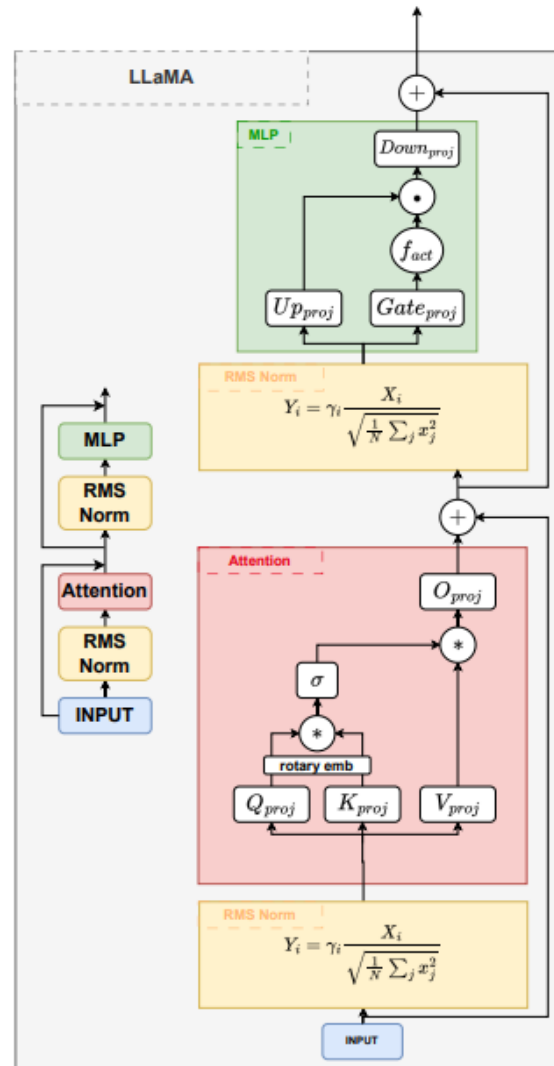
Faster-RCNN applied to ADAS:



Transformers: ViT and LLM



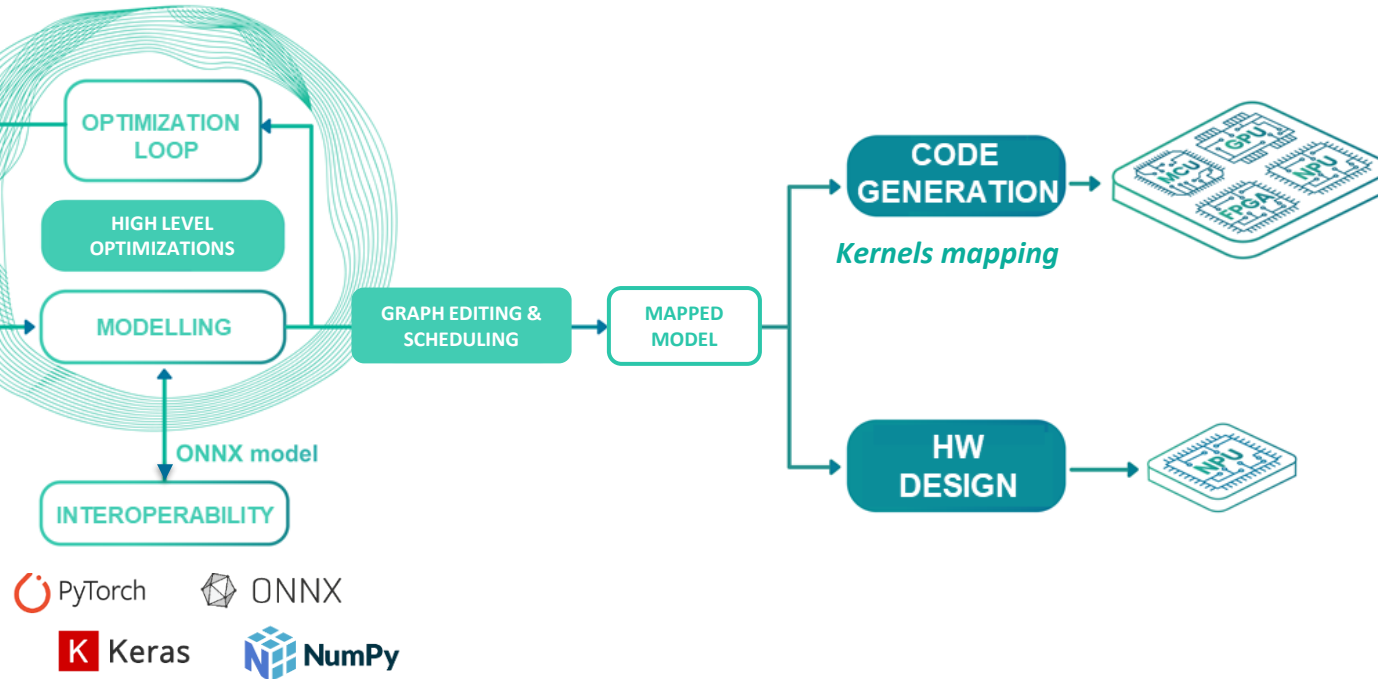
(a) OPT



(b) LLaMA

- LLM are causal models
- They only predict the next token
- 3 main blocs:
 - Attention
 - Feed-Forward Network
 - Normalisation

Aidge: a DL framework dedicated to embedded AI



Remember: the targeted hardware does not necessarily have a compiler!

End-to-end interoperable framework including optimization and deployment

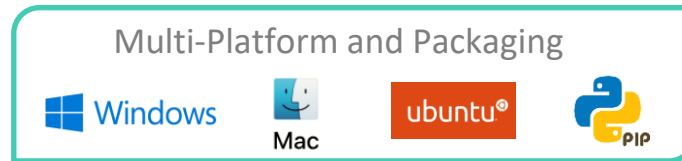
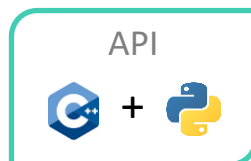
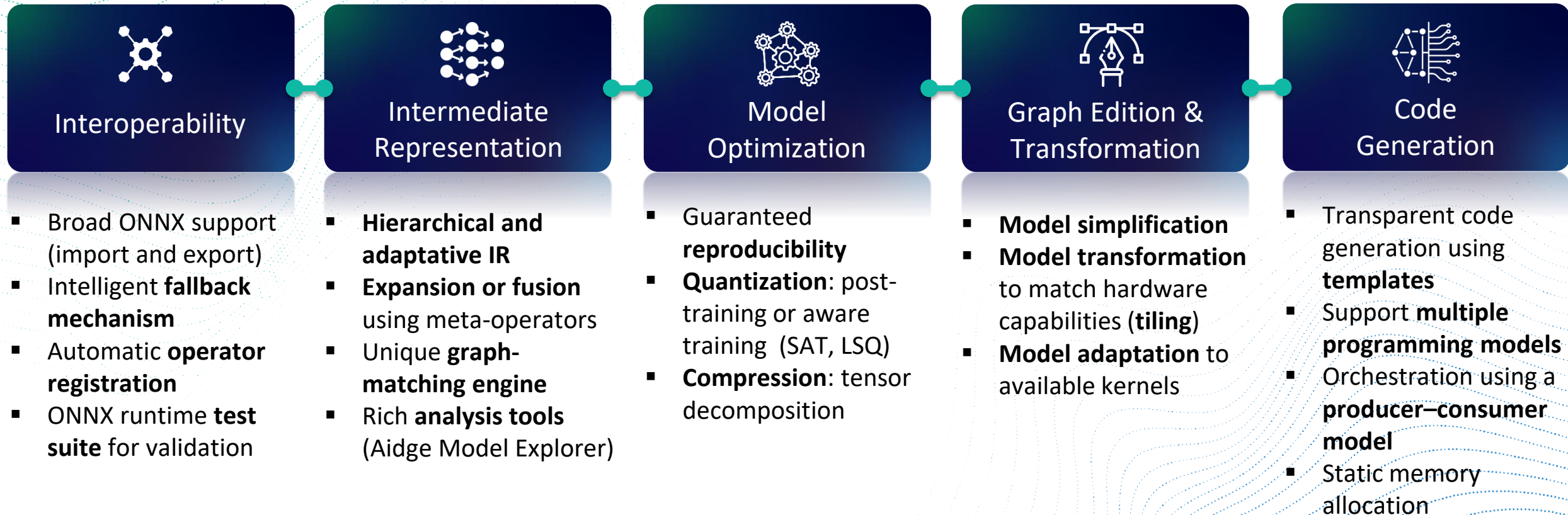
Modular library with light core, plugins system and few dependencies

Open source code hosted by **ECLIPSE FOUNDATION** with adapted license to create proprietary software (EPLv2)

+40 industrial and academics partners



Modular deployment framework: beyond fragmentation & black-box



The right IR is everything

Graph matching and rewrite is **the single most important operation** for optimizing and mapping high level algorithms to hardware compute resources:

- Strong ONNX interoperability require graph adaptation for **import and export**
- **ONNX simplifiers** consist of matching sub-graph pattern and replacing them
- **Quantization** consists of inserting quantizer nodes, scaling nodes and performing nodes fusion
- Tensor-decomposition **compression** consists of replacing Conv with several smaller Conv
- Mapping to memory-limited physical computing resources require operators **tiling**
- Mapping to available compute kernels require **graph adaptation** (casting, transposition, packing...)
- Multi-modal optimization with data augmentation require synchronized data processing pipelines

... all these can be described extremely efficiently with the **right intermediate representation (IR)**!

Key features of Aidge's IR

Key features:

A transparent representation with a graph

A multi-paradigm abstraction model

A hierarchical graph model

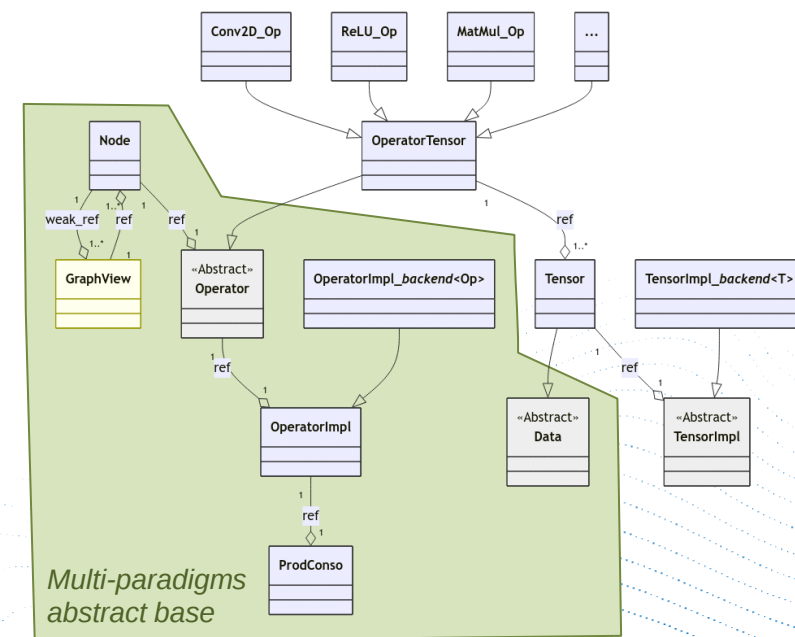
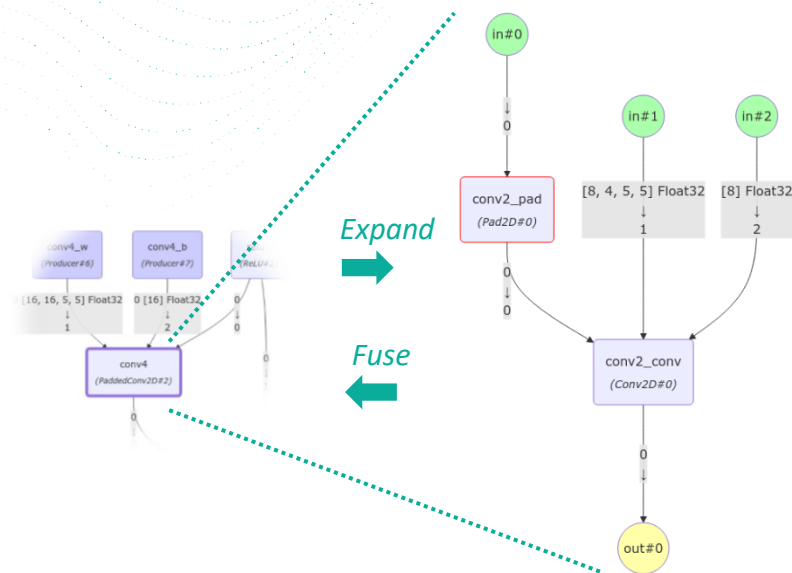
A generic graph model

An accessible graph model

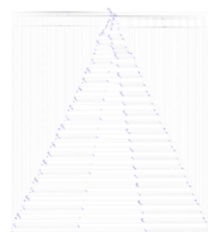
→ Meta Operator

→ Generic Operator

→ Graph Matching



```
dinov2_model = aidge_onnx.load_onnx("dinov2.onnx")
aidge_core.fuse_to_metaops(dinov2_model, "MatMul->Add", "Linear")
aidge_core.fuse_to_metaops(dinov2_model, "MatMul->Div#1->Softmax->MatMul;"
                                "Div#1<1~Producer", "ScaledDotProductAttention")
aidge_core.fuse_to_metaops(dinov2_model, "ScaledDotProductAttention#1->Transpose->Reshape#1->Linear;"
                                "Reshape#1<1~Producer;"
                                "ScaledDotProductAttention#1<0-(Transpose<-Reshape#2<-Add#1);"
                                "ScaledDotProductAttention#1<1-(Transpose<-Reshape#3<-Add#2);"
                                "ScaledDotProductAttention#1<2-(Transpose<-Reshape#4<-Add#3);"
                                "Reshape#2<1~Producer;"
                                "Add#1<*-0-Split#1;"
                                "Add#2<*-1-Split#1;"
                                "Add#3<*-2-Split#1;"
                                "Split#1<-MatMul;"
                                "Split#1<1~Producer", "MultiHeadAttention")
```



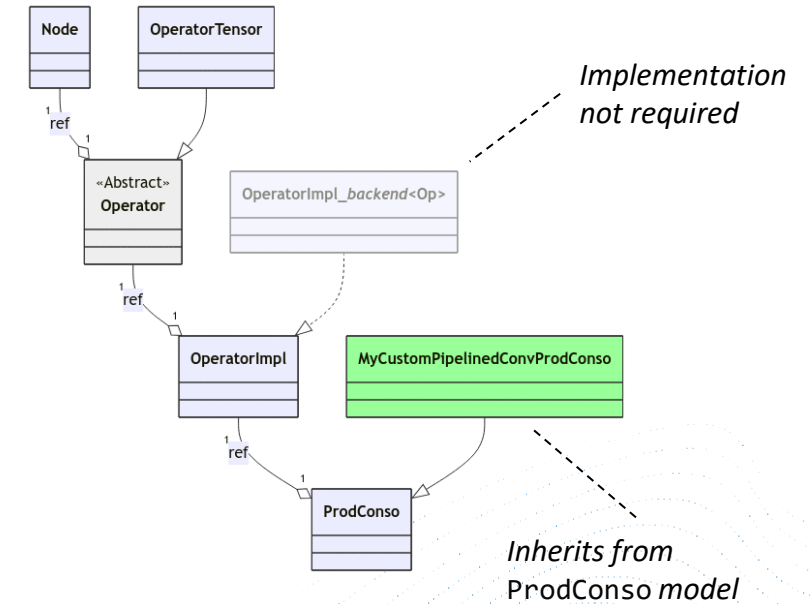
Key features of Aideg's IR

- Aideg introduces a well-defined consumer-producer (C-P) model
 - Similar to transaction-level modeling (TLM) for electronic design

C-P model is
state full

```
class Conv2D_DataFlow_CP(aideg_core.ProdConso):  
    def __init__(self, op: aideg_core.Operator):  
        aideg_core.ProdConso.__init__(self, op)  
        self.state_begin = True  
  
    def get_nb_required_data(self, input_idx):  
        input = self.get_operator().get_input(input_idx)  
        if input:  
            if self.state_begin:  
                # Require 3 lines for 3x3 kernel  
                return aideg_core.Elts_t.data_elts(3 * input.dims()[2])  
            else:  
                return aideg_core.Elts_t.data_elts(input.dims()[2])  
        else:  
            return aideg_core.Elts_t.none_elts()  
  
    def get_required_memory(self, output_idx, inputs_size):  
        output = self.get_operator().get_output(output_idx)  
        self.state_begin = False  
        return aideg_core.Elts_t.data_elts(output.dims()[2])
```

Fine grain (pipelined)
dataflow modeling



C-P model	Specification
Token-based	Arbitrary data quantity (FSM logic / Petri net modeling) <i>Standard Tensor-based AI frameworks scheduling</i>
Data-based	Precise amount of data elements consumed and produced at each execution step
Hybrid	Data-based to token-based

Key features of Aidege's IR

- Aidege's scheduling is **static**
 - Get early and late logical static scheduling for any graph, cyclic or acyclic

```
lstm = aidege_core.LSTM(in_channels=4, hidden_channels=8, seq_length=5)
```

```
# Flatten the graph:
```

```
lstm_model = aidege_core.get_connected_graph_view(lstm)  
aidege_core.expand_metaops(lstm_model)  
lstm_model.set_backend("cpu")
```

```
# Create the Scheduler
```

```
lstm_scheduler = aidege_core.SequentialScheduler(lstm_model)  
lstm_scheduler.generate_scheduling()
```

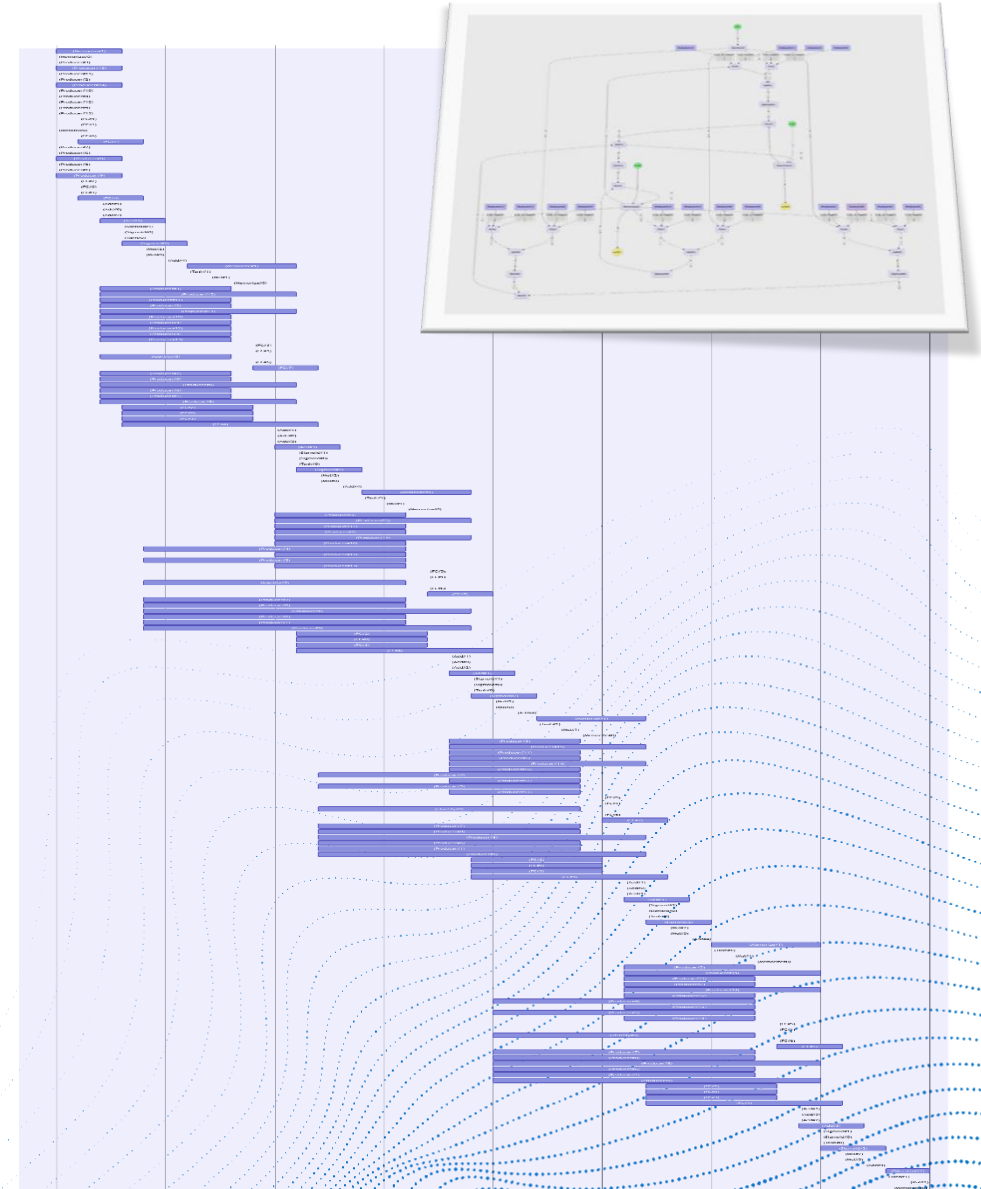
```
# Display static scheduling
```

```
lstm_scheduler.save_static_scheduling_diagram("lstm_scheduling")
```

Based on the specified C-P model

Agnostic of the programming model

*An associated operator implementation is not required
(the scheduling is decoupled from the execution path)*



Key features of Aideg's IR

➤ Comprehensive and simple static analysis

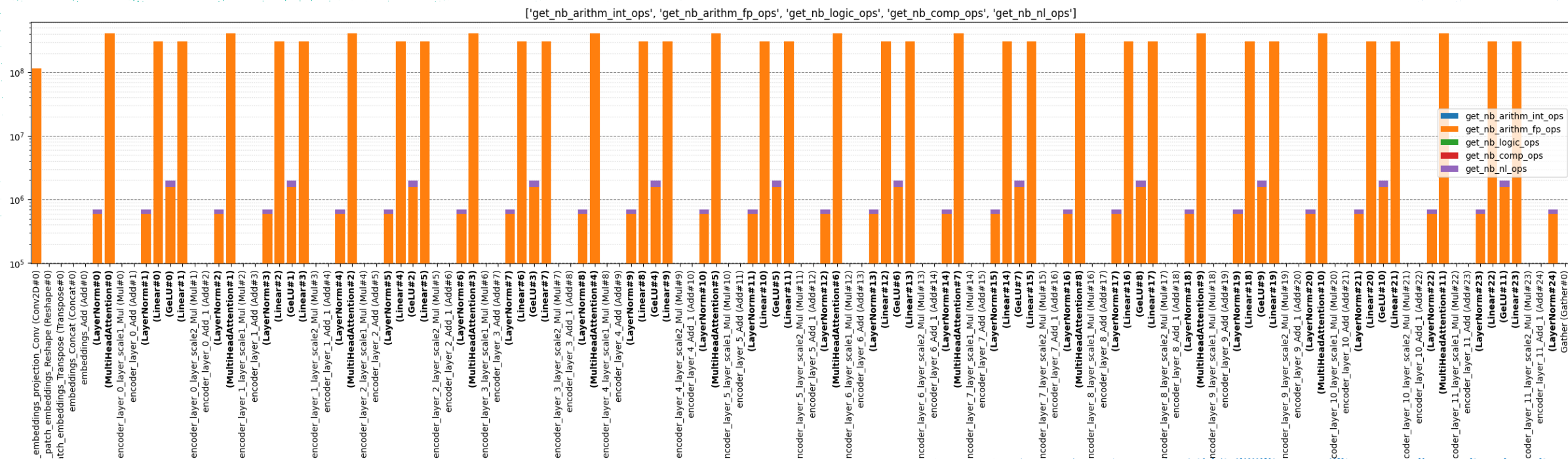
```
stats = aideg_core.static_analysis.StaticAnalysisExt(dinov2_model)
stats.summary() # Simple text output similar to Keras model.summary()

stats.log_nb_ops_by_type("stats_ops.png", log_scale=True)
```

Rich low-level & high-level APIs:

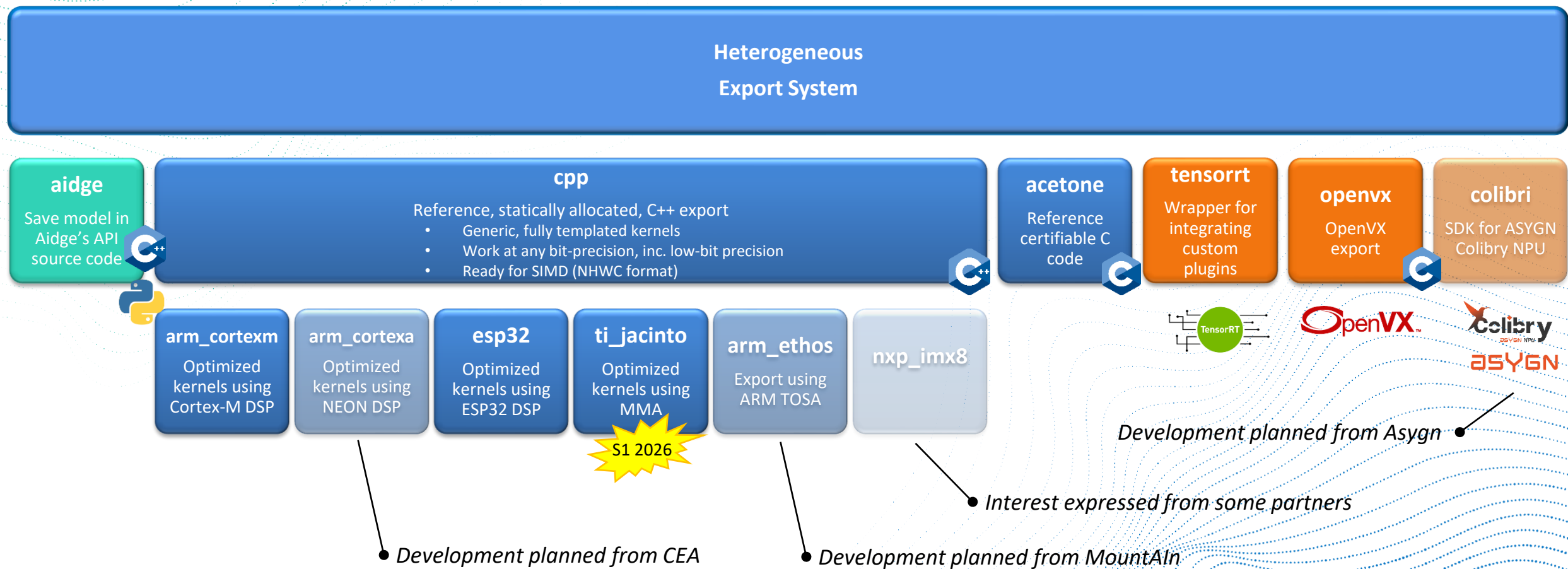
- Base C++: for lower-level heuristics integration
- Extended Python: with additional visualization methods

Auto aggregation for Meta operators,
specifiable for Generic operators



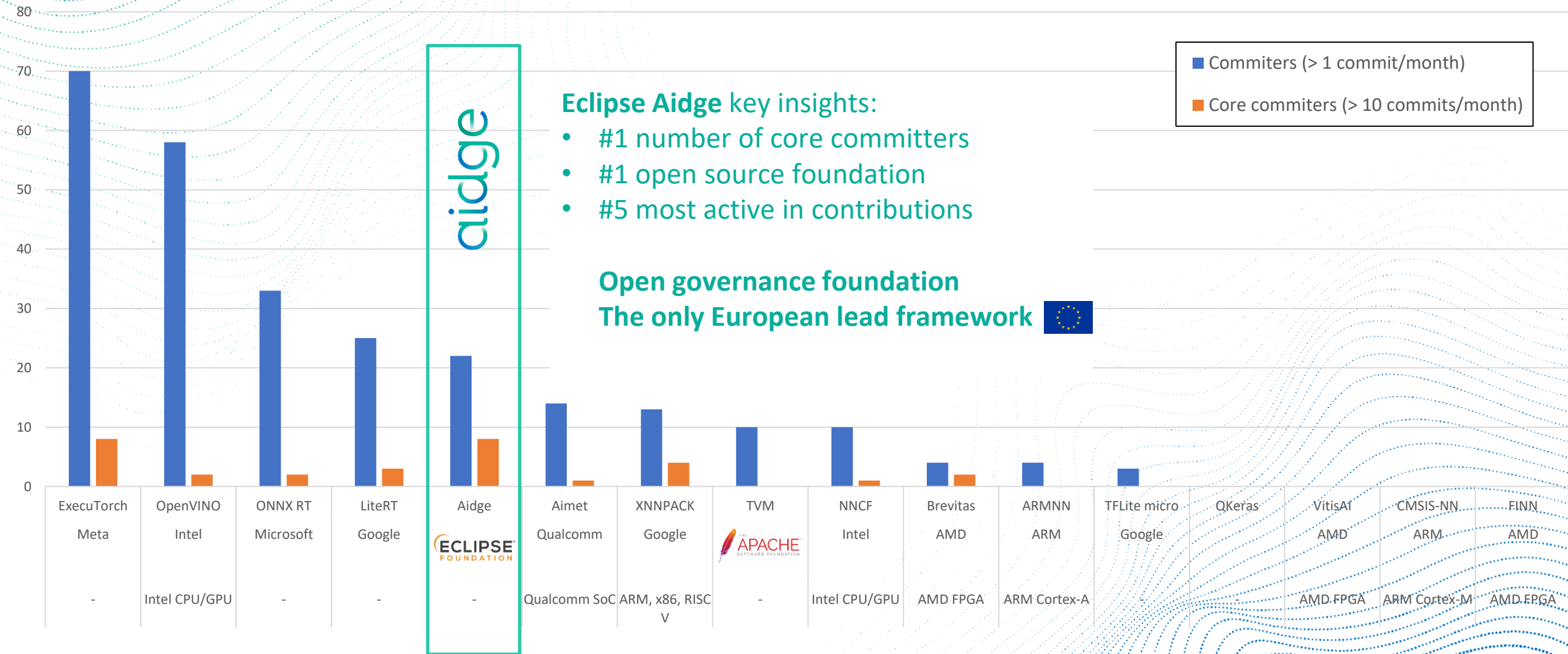
Aidge's export modules overview

Code generation exports



Aidge in the open source landscape

Project activity in the last 12 months (sept. 24 – sept. 25)



AI Circuits & Architectures

Transformers & multimodality

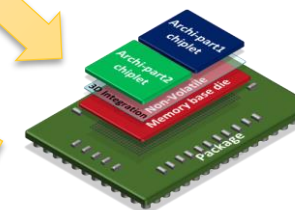


BUMBLEBEE
Attention based AI
Circuit Architecture

CNN backbone

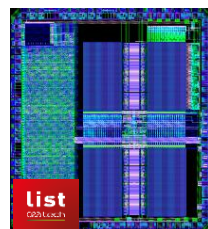


NeuroCORG1
Feature Extractor (HD images)
GF22nm, OxRAM



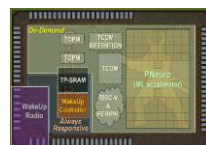
MOSAIC
Modular AI systems composed
of heterogeneous components

Smart Sensors



Esperanto
Gesture detection
130nm + OxRAM + PMUT,

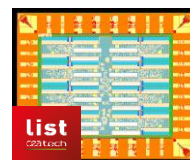
CNN



PNEURO
Raptor HW/SW IP
FD28nm

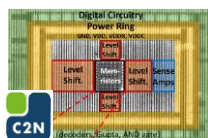
DOLPHIN
DESIGN

Spiking SNN

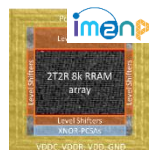


LARGO
28nm, OxRAM

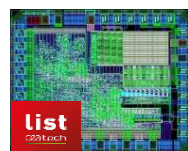
Classification



Bayesian
machine



BNN
Digital popcount
130nm, OxRAM



SPIRIT
Spiking SNN
130nm, OxRAM

2018

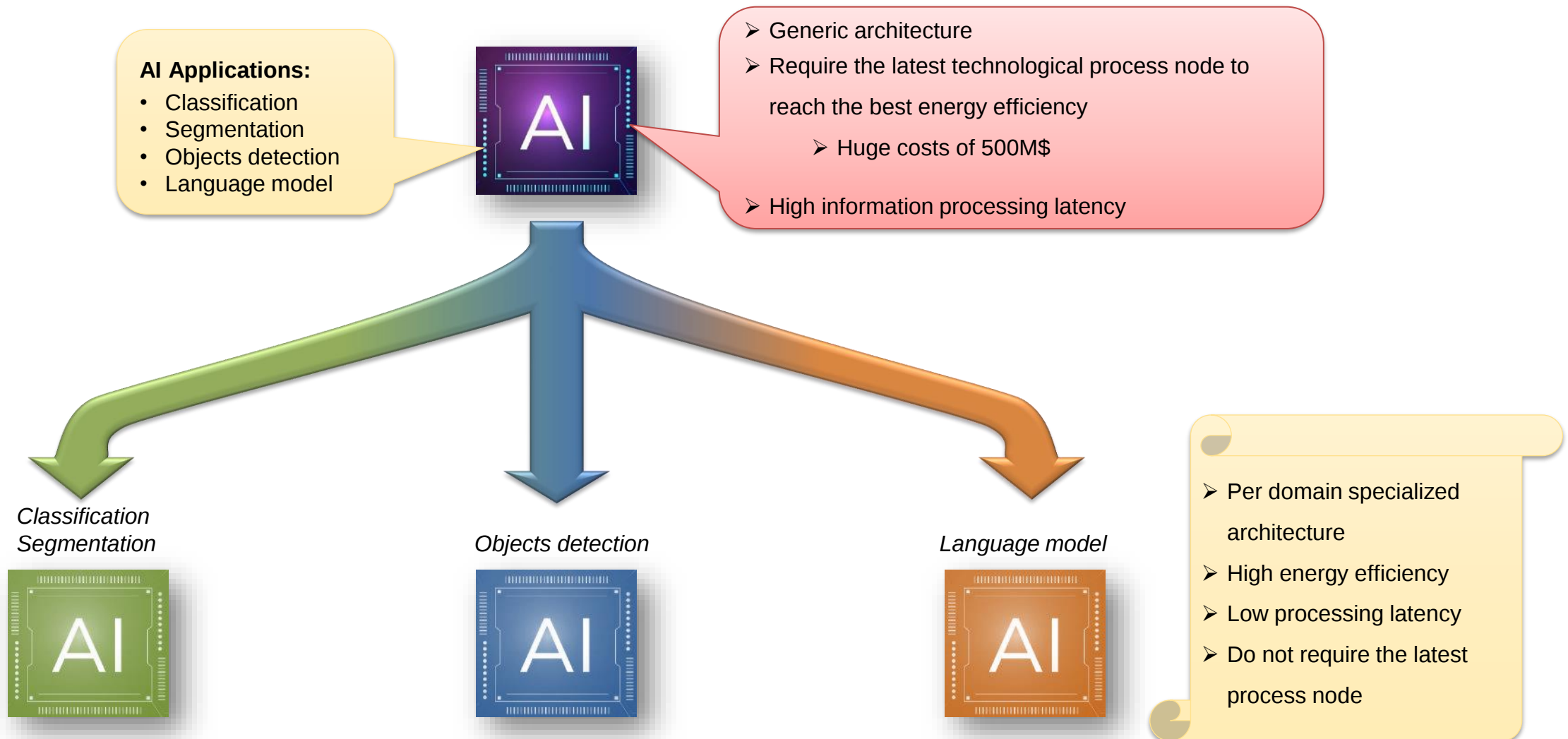
2020

2022

2023

2026

Towards hardware *generation*

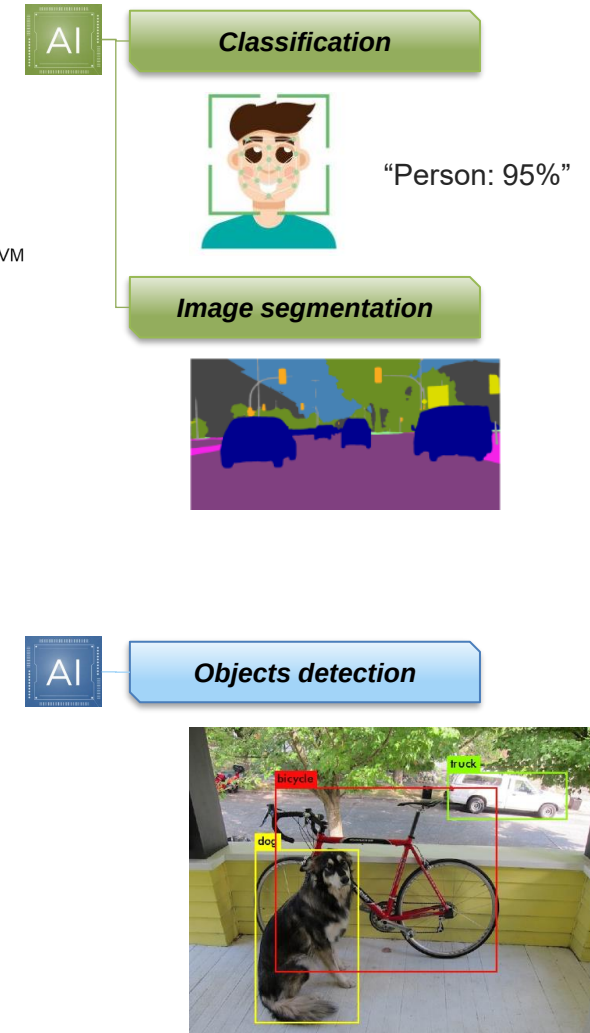
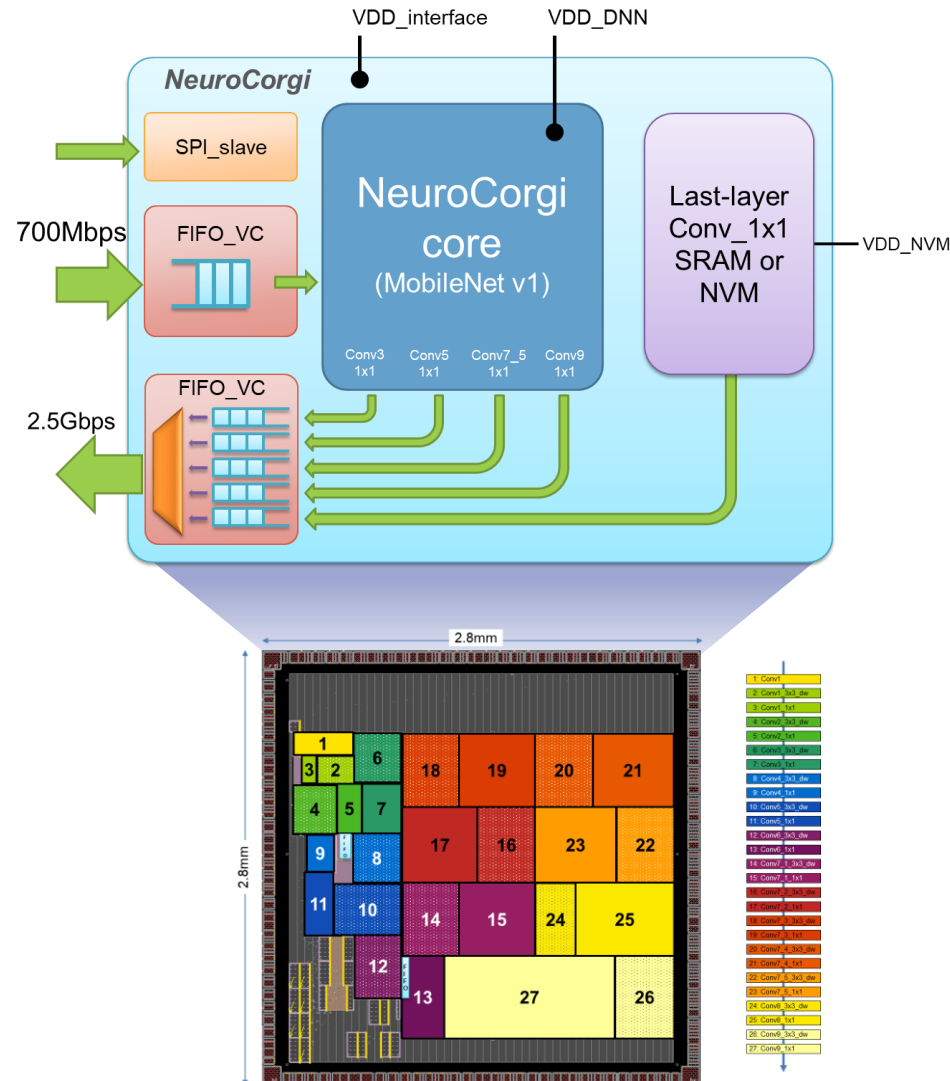


The NeuroCORGI circuit

- Fix topology: MobileNet v1
- Parameters freeze at design
- 4-bits quantization
- Very low clock frequency
- NVM-Leti memory
- 3 different circuit versions:



⇒ Ultra low energy consumption
 ⇒ Ultra low processing latency



NeuroCORGI benchmark results

Ultra low power: Active = **26mW** → **26%** of imager power (>100mW) ; Leakage = **218μW**

Best in class energy / inference / pixel: **0,84nJ** → **8x** to **1000x** better compared to SoTA

Best in class latency with batch size 1: **14ms** → **42%** of 30fps imager latency

	This work	Samsung (ISSCC 2022)	Vega (ISSCC 2021)	MediaTek (ISSCC2020)	Alibaba (ISSCC 2020)	KAIST	Eyeriss v2	Dolphin RAPTOR	GAP8
Technology	GF22FDX	4nm	GF22FDX	7nm	TSMC 12nm	65nm	TSMC 65nm	GF22FDX	
AI model	MobileNet v1	MobileNet TPU	MobileNet v2	MobileNet v1	ResNet v1	MobileNet	MobileNet v1 (alpha 0.5)	MobileNet v2	MobileNet v1
Energy/inference (224x224)	0.042mJ	0.35mJ	1.19mJ	0.17mJ	2.00 – 3.51 mJ	1.2mJ	1.55mJ*	2.6mJ	49.7mJ
Energy/inference/ pixel	0.84nJ	6.90nJ	23.72nJ	3.46nJ	39.92 – 70.02 nJ	23.92nJ	30.99nJ	51.82nJ	990.05nJ
Energy/inference/ pixel W.r.t. this work	1x	8.19x	28.15x	4.11x	47.31x – 83.10x	36.78x	36.78x	61.5x	1175x

Approach proven on several applications



Drones/USV



**Underwater
Acoustic Signal
Classification**



**Robust
Autonomous
Landing**



**3D Object Detection
and Classification
of Road Users**



**Autonomous
Weeding System**



**Tomato pests and
diseases forecast**



list



Olivier Bichler