



The background features a light blue hexagonal grid. In the upper left, a cluster of small white triangles forms a larger hexagonal shape. Scattered across the scene are numerous 3D cubes of various sizes and colors, including red, orange, yellow, and blue. Some cubes have a bright white highlight on one of their faces. Faint mathematical formulas are visible on some of the background cubes, such as $F(\eta) = e^{-\eta} \left(\frac{\alpha}{1+\alpha} \right)$, $f_{\alpha}(0) = 1 - e^{-\eta}$, and $F(\eta) = e^{-\eta}$.

Improving Chip Design Performance and Productivity Using Machine Learning

Narender Hanchate

Sr Machine Learning Architect, Digital and Signoff Group, Cadence

 **cā dence**[®]

Today's Chip Design Market Landscape

The semiconductor industry is experiencing a renaissance

- Strong growth in 5G, autonomous driving, hyperscale compute, and industrial IoT
- Underlying each of these trends is the application of artificial intelligence (AI) and machine learning (ML)

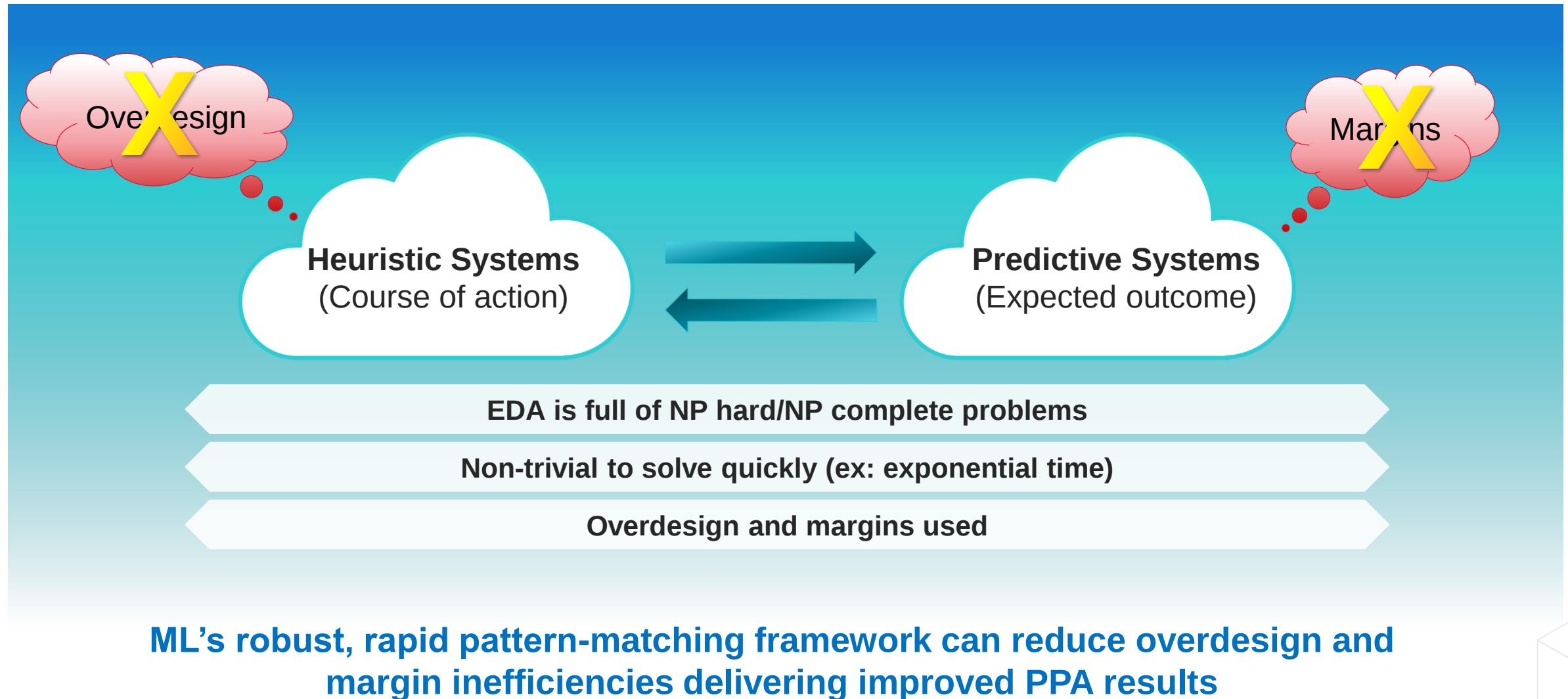
New applications and technological interdependencies are generating demand for even more compute, more functionality, faster data transmission speeds

- Today's electronics have more chips in them with no end to this trend in sight

Next-generation chips must be produced faster and smarter

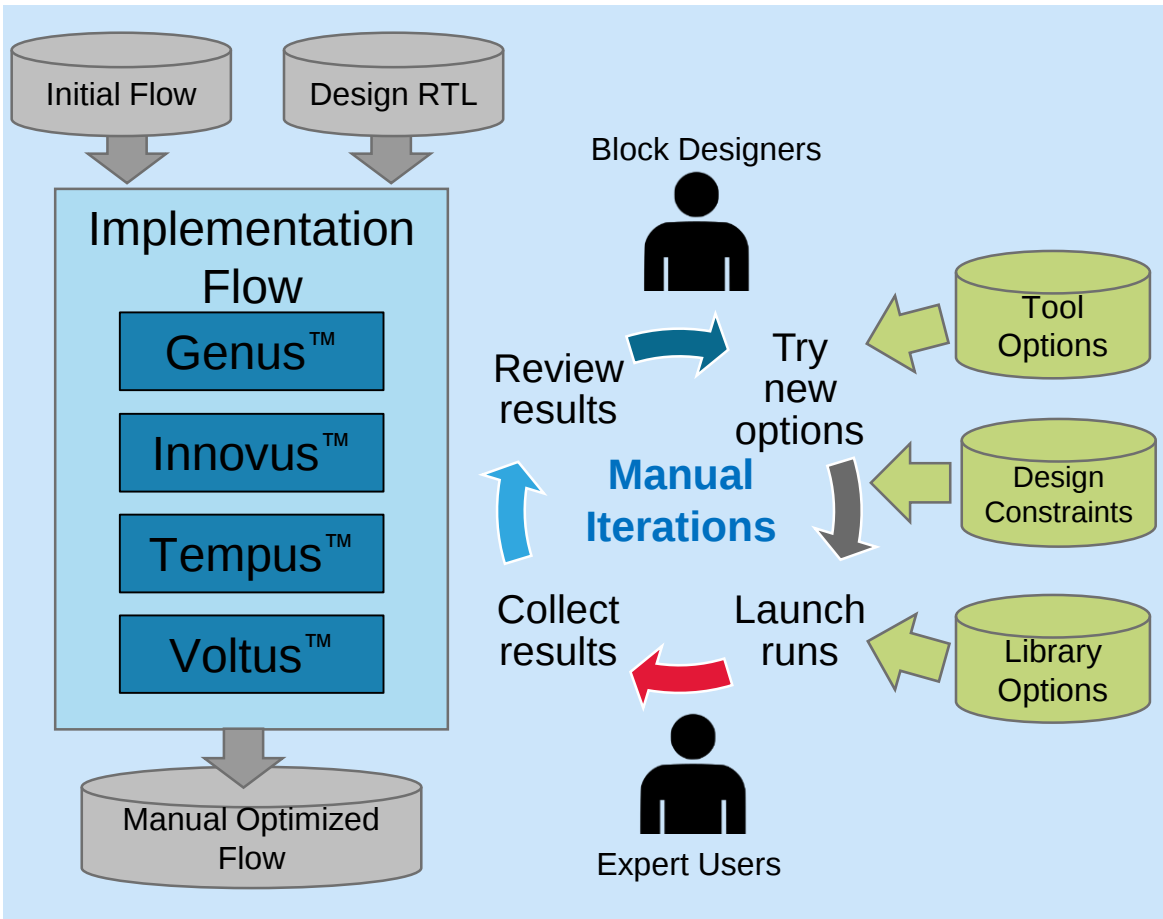
- Engineers are overloaded and need support to keep up with demand

Machine Learning Is a Good Fit for Digital Implementation



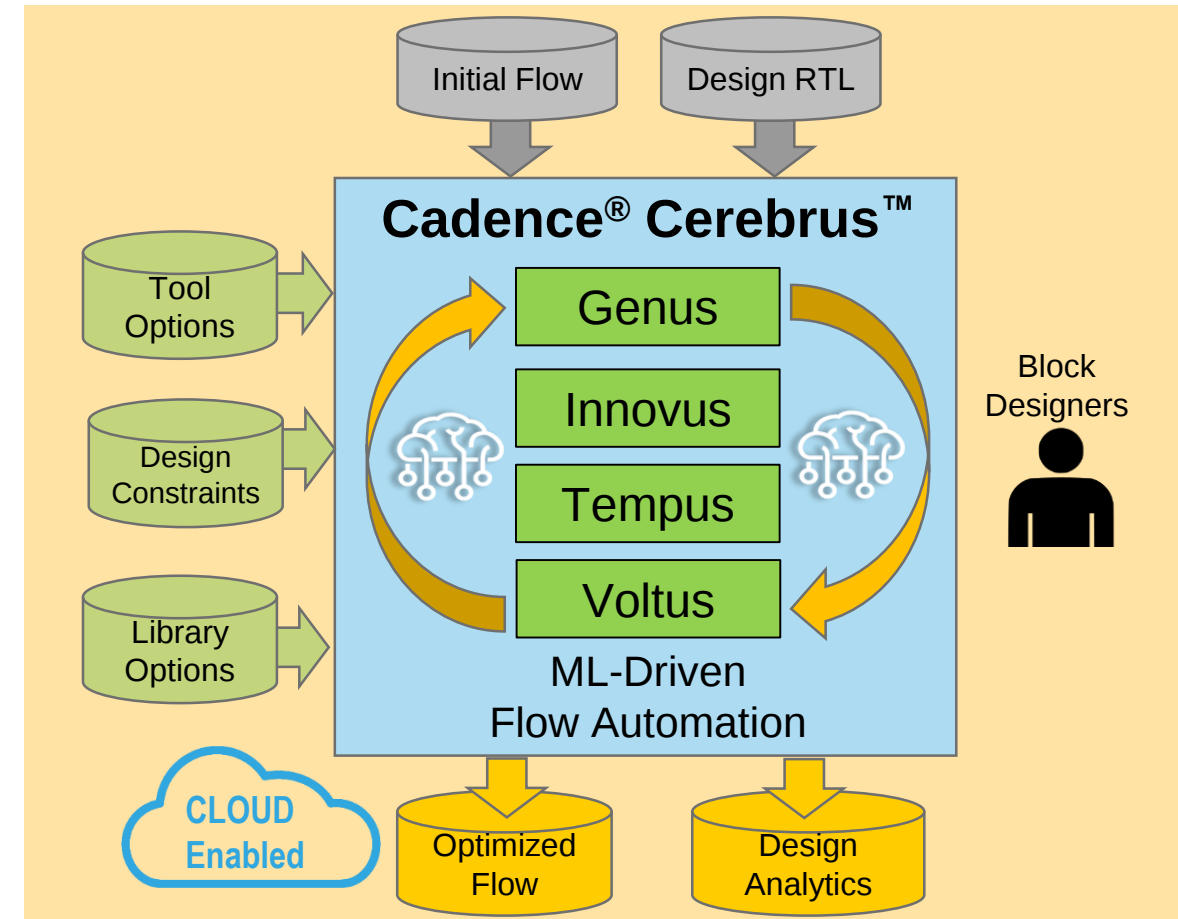
Cadence Cerebrus

Traditional Flow



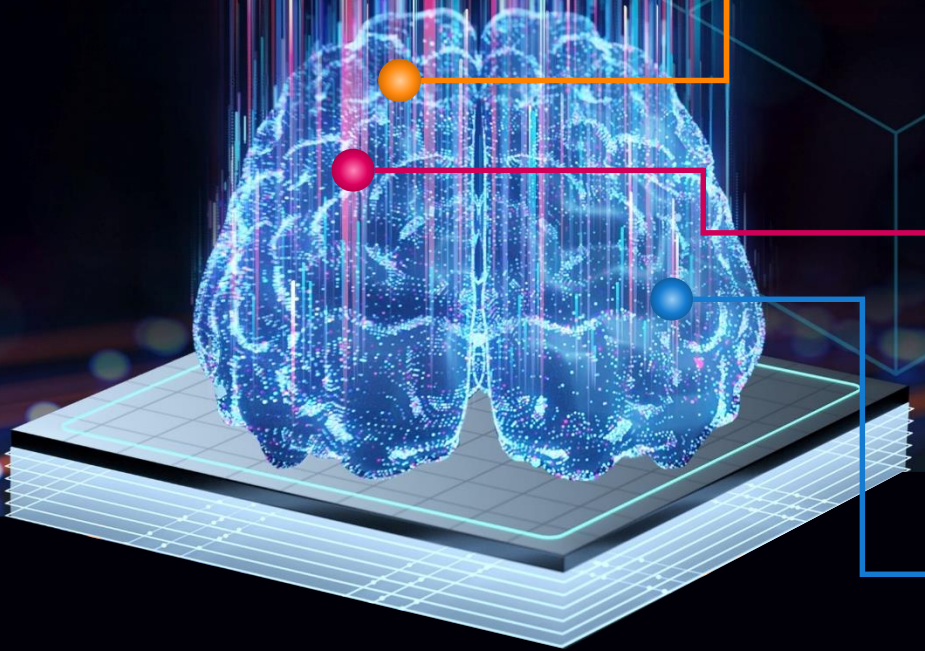
- High engineering effort
- Unpredictable schedules

Cadence Cerebrus ML Flow



- Improved designer productivity
- Better PPA
- Efficient compute usage

Cadence Cerebrus: The Future of Intelligent Chip Design



New ML-based tool that
automates and scales
digital chip design

Productivity and PPA Revolution

- Unique reinforcement ML
- Delivers up to 10X better productivity and 20% PPA improvements

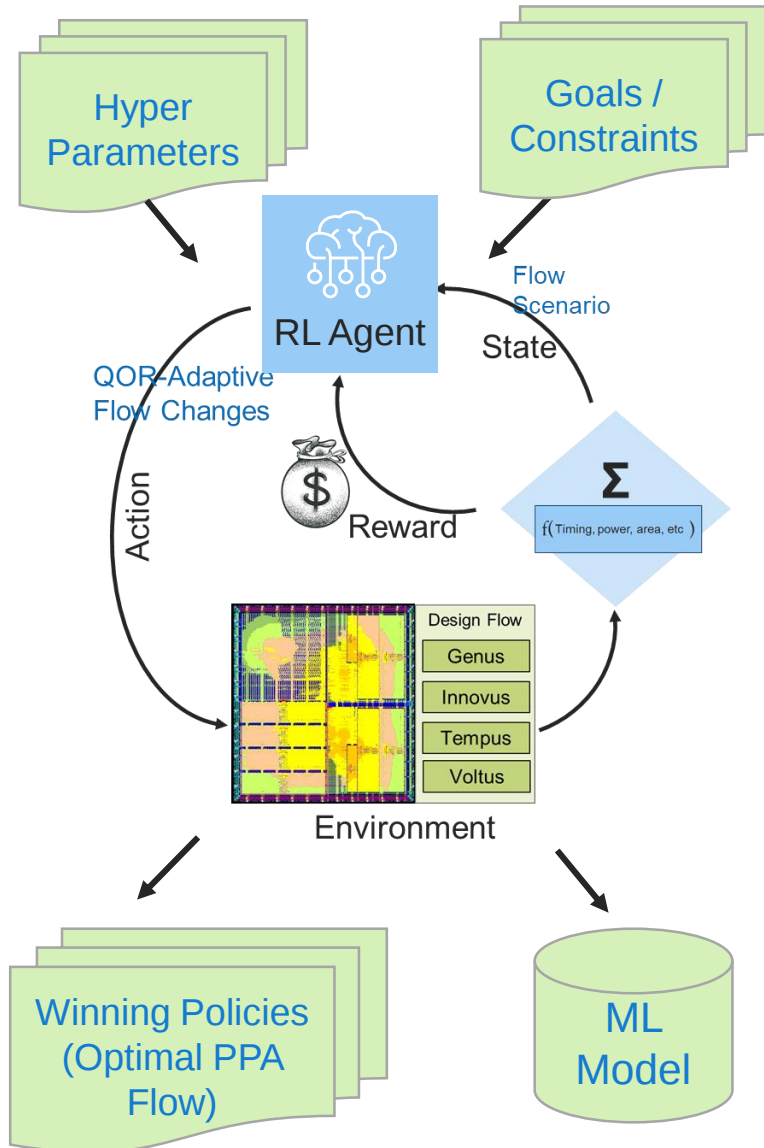
Automated RTL to GDS Full Flow Optimization

- Delivers better PPA more quickly
- Improves engineering team productivity

Scalable, Distributed Computing Solution

- On-premises or cloud computing resources
- Efficient, scalable solution as design size and complexity grow

Applying Reinforcement Learning Principles



- RL requires efficient exploration
- Scenario creation driven by probability of success
- Smart decision points (Early Stop / Reuse etc.)
- Resource management (CPU / Disk)

Cadence Cerebrus – Full Flow Design Optimization

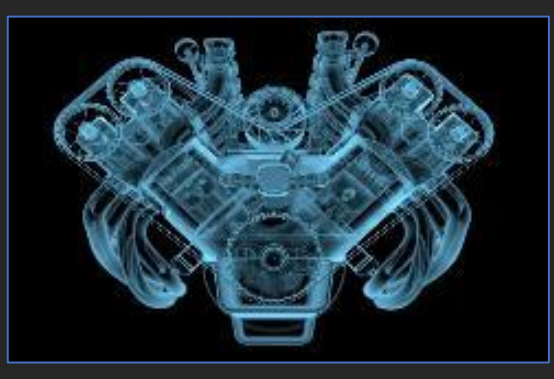
Design
Stratus™, Modus,
Conformal®

Implementation
Genus™, Innovus™,
Joules™

Electrical Signoff
Quantus™, Tempus™
Voltus™, Liberate™

Physical Signoff
Pegasus™, DFM

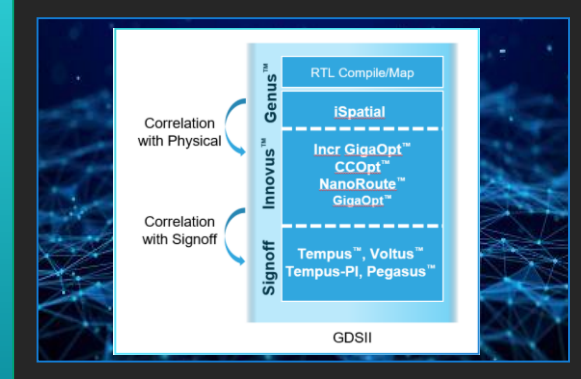
Cadence Digital Full Flow
RTL – GDSII



- Best-in-class optimization engine
- Signoff accuracy
- Single CPU performance
- Efficient memory management

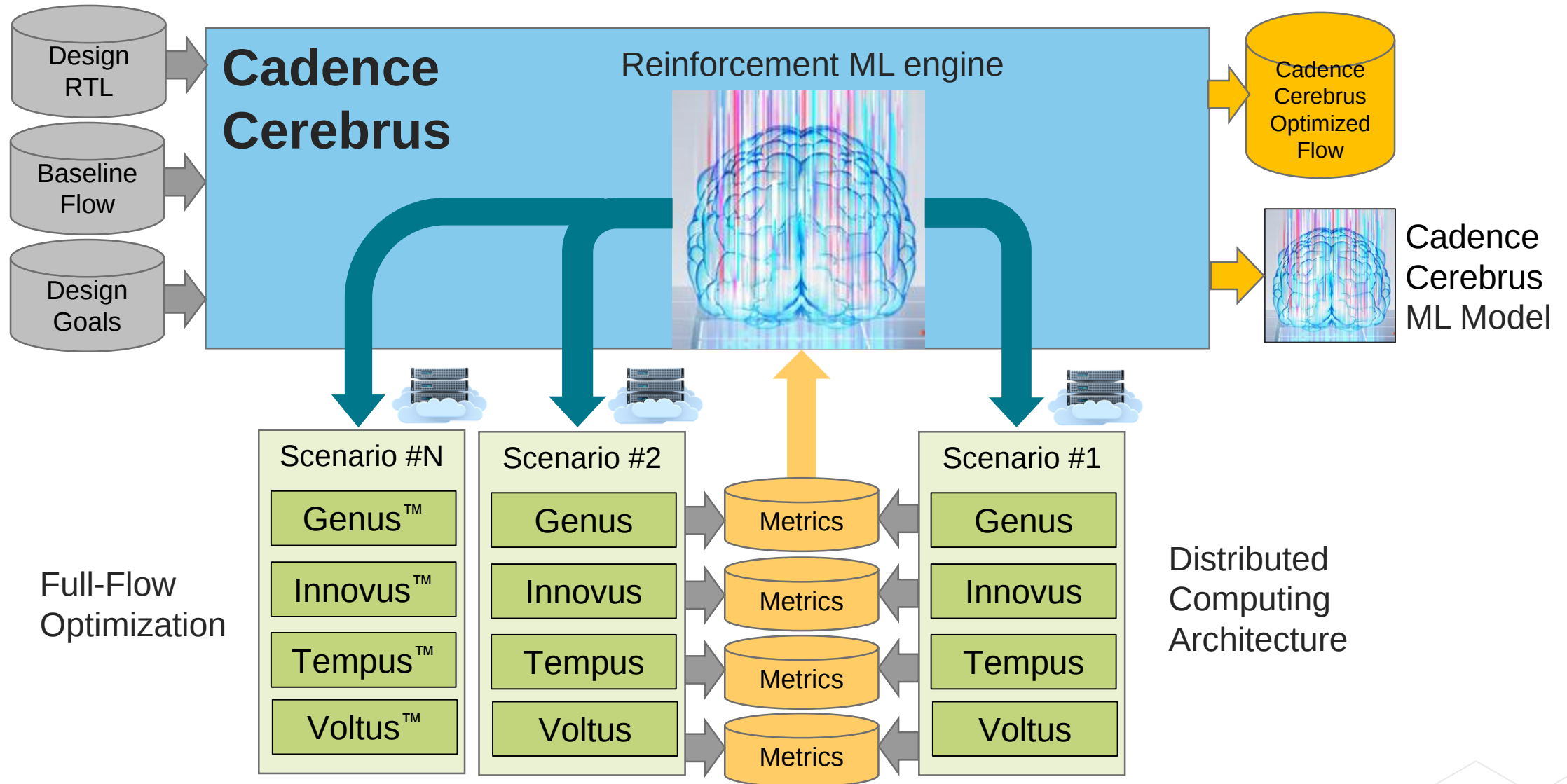


- Massively parallel
 - Multi-core and multi-thread
 - Fully distributed
- Cloud ready
- Best-in-class hierarchical solution



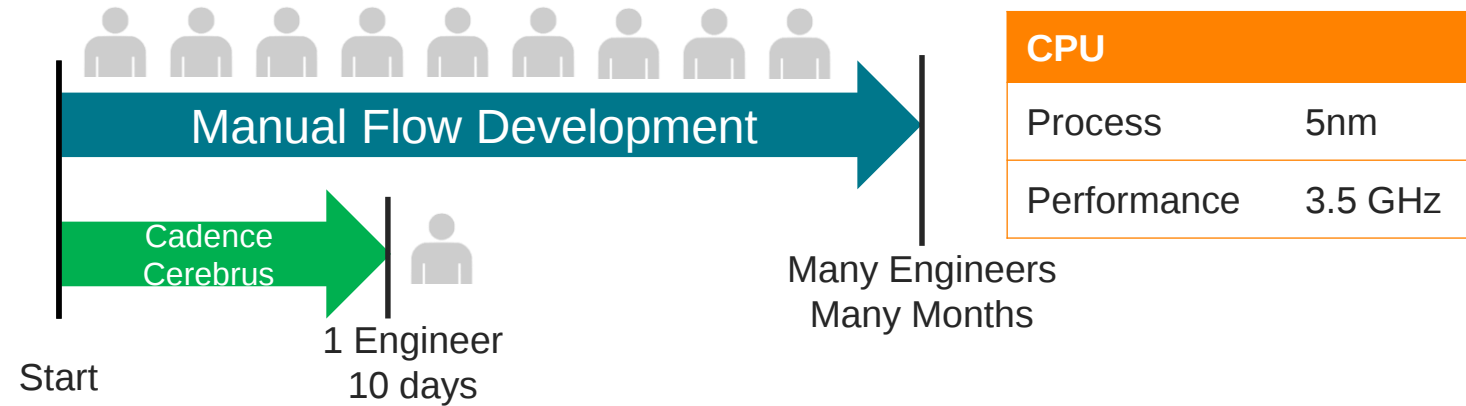
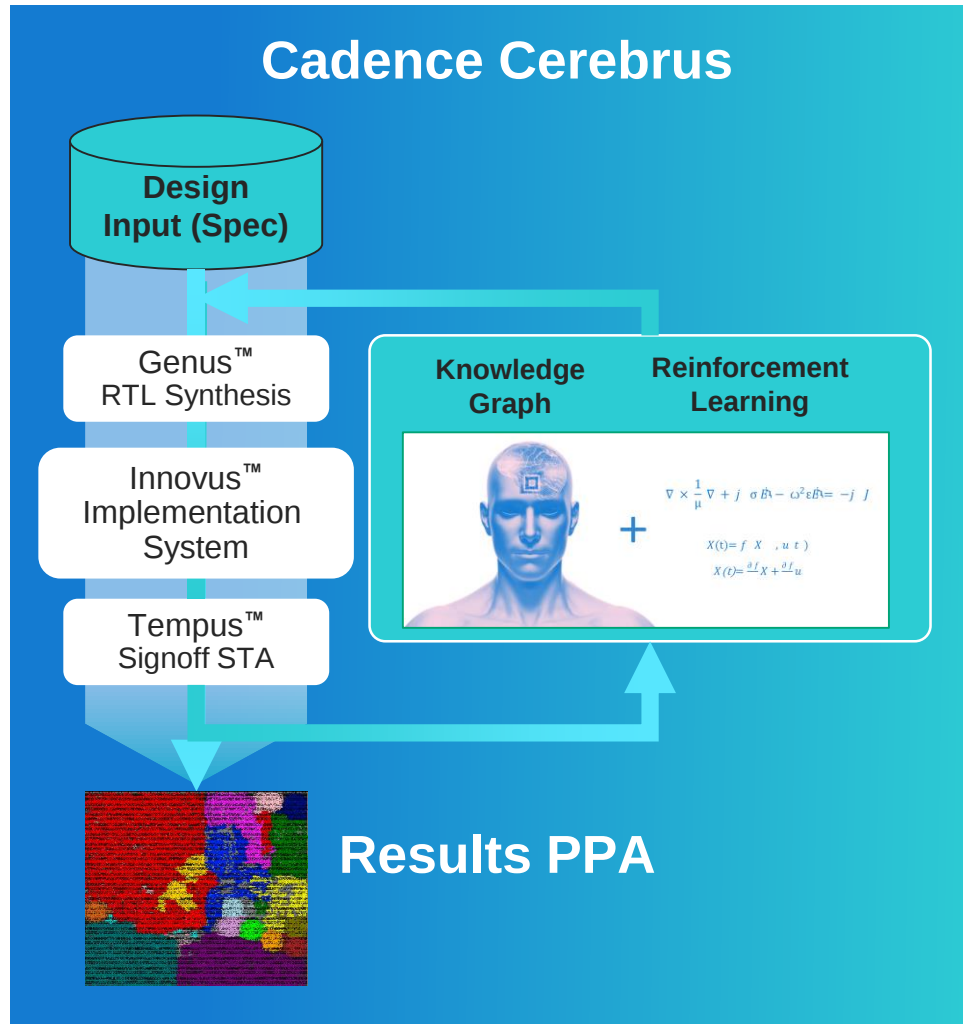
- Digital full-flow PPA excellence
- Mixed signal
- Natively shared engines
- Machine learning automation

Cadence Cerebrus - Distributed Computing & ML Architecture



Cadence Cerebrus - ML for Better PPA & Full Flow Productivity

ML flow optimization improves designer efficiency

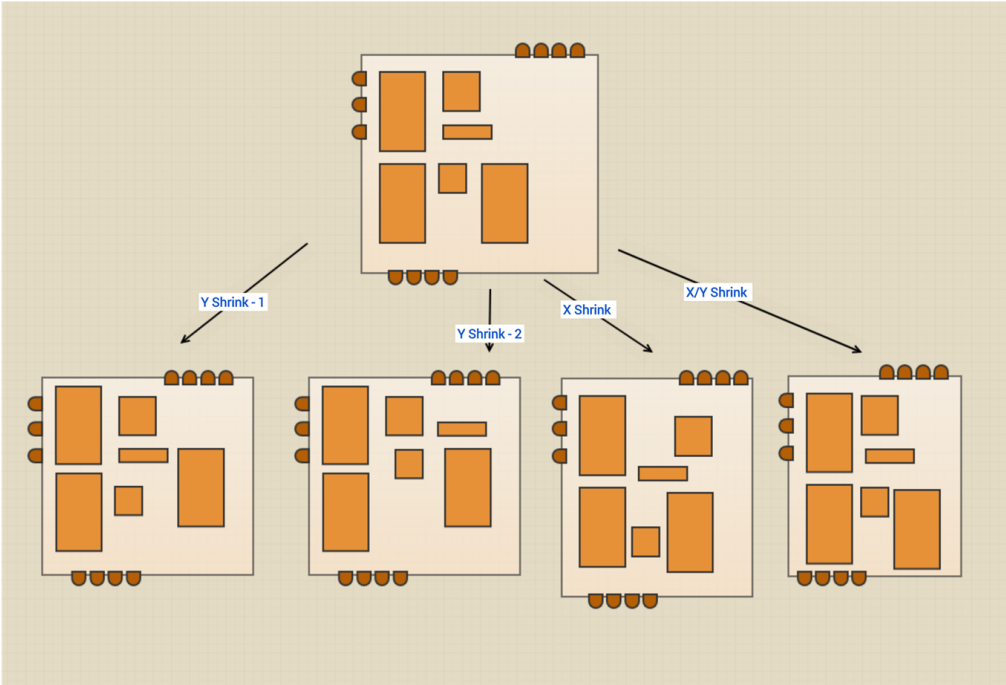


- Cadence® Cerebrus™ automatically improved PPA of 5nm mobile CPU
- Cadence Cerebrus used 30 parallel jobs
- Within 10 days converged on improved flow

Cadence Cerebrus Improvements over Baseline

Parameter	Improvement	Percent
Performance	420 MHz	14%
Leakage power	26 mW	7%
Total power	62 mW	3%
Density		5%

Cadence Cerebrus – Automated Floorplan Optimization



- Floorplan can be automatically resized in any direction
- Innovus mixed placer used to find optimal macro location in resized floorplan

Design – CPU Core	
Process	12nm
Performance	2 GHz

- Customer wanted to achieve 2GHz on latest CPU implementation
- Cadence® Cerebrus™ ran 50 flow experiments
 - Used Innovus™ mixed placer to improve floorplan
 - Other flow changes delivered better performance

Cadence Cerebrus Improvements over Baseline

Parameter	Improvement
Performance	+200MHz
Total failing timing	83%
Leakage power	17%

Cadence Cerebrus – ML Capabilities Improved Flow, PPA, and Productivity

“

To efficiently maximize the performance of new products that use emerging process nodes, digital implementation flows used by our engineering team need to be continuously updated. Automated design flow optimization is critical for realizing product development at a much higher throughput. Cadence Cerebrus, with its innovative ML capabilities, and the Cadence RTL-to-signoff tools have provided automated flow optimization and floorplan exploration, improving design performance by more than 10%. Following this success, the new approach will be adopted in the development of our latest design projects.

Satoshi Shibatani,

*Director, Digital Design Technology Department,
Shared R&D EDA Division, Renesas*

As Samsung Foundry continues to deploy up-to-date process nodes, the efficiency of our Design Technology Co-Optimization (DTCO) program is very important, and we are always looking for innovative ways to exceed PPA in chip implementation. As part of our long-term partnership with Cadence, Samsung Foundry has used Cadence Cerebrus and the Cadence digital implementation flow on multiple applications. We've observed more than an 8% power reduction on some of our most critical blocks in just a few days versus many months of manual effort. In addition, we are using Cerebrus for automated floorplan power distribution network sizing, which has resulted in more than 50% better final design timing. Due to Cadence Cerebrus and the digital implementation flow delivering better PPA and significant productivity improvements, the solution has become a valuable addition to our DTCO program.

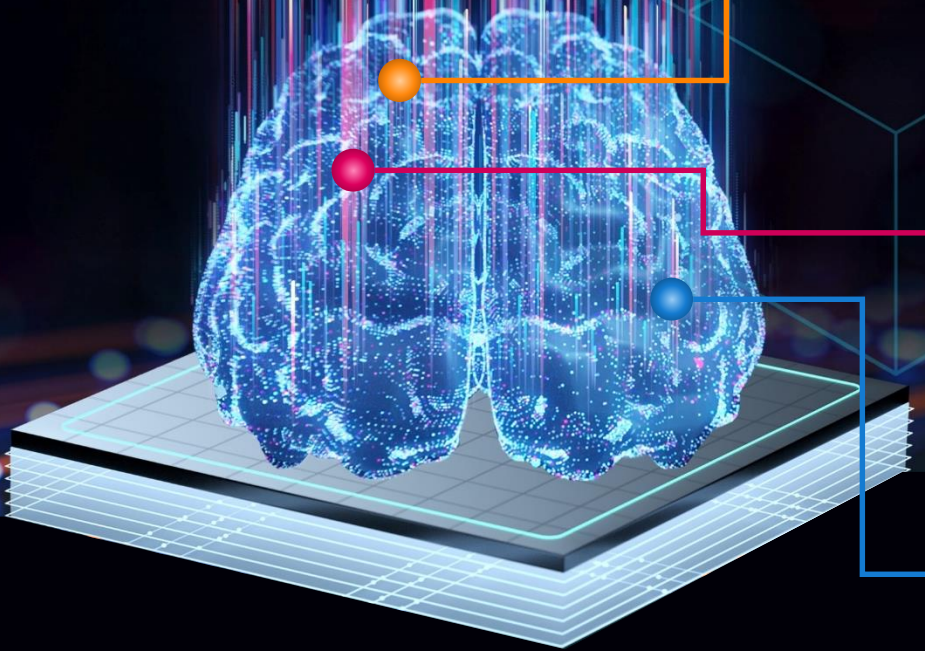
Sangyun Kim,

Vice President, Design Technology, Samsung Foundry

”

cadence

Cadence Cerebrus: The Future of Intelligent Chip Design



New ML-based tool that
automates and scales
digital chip design

Productivity and PPA Revolution

- Unique reinforcement ML
- Delivers up to 10X better productivity and 20% PPA improvements

Automated RTL to GDS Full-Flow Optimization

- Delivers better PPA more quickly
- Improves engineering team productivity

Scalable, Distributed Computing Solution

- On-premises or cloud computing resources
- Efficient, scalable solution as design size and complexity grow



Find more Cadence® Cerebrus™ material at

www.cadence.com/go/cerebrus

- White paper
 - Videos
- Product Information

https://www.cadence.com/en_US/home/multimedia.html/content/dam/cadence-www/global/en_US/videos/solutions/renesas-designed-with-cadence.mp4

https://www.cadence.com/en_US/home/multimedia.html/content/dam/cadence-www/global/en_US/videos/solutions/samsung-designed-with-cadence.mp4



cādence®

© 2022 Cadence Design Systems, Inc. All rights reserved worldwide. Cadence, the Cadence logo, and the other Cadence marks found at <https://www.cadence.com/go/trademarks> are trademarks or registered trademarks of Cadence Design Systems, Inc. Accellera and SystemC are trademarks of Accellera Systems Initiative Inc. All Arm products are registered trademarks or trademarks of Arm Limited (or its subsidiaries) in the US and/or elsewhere. All MIPI specifications are registered trademarks or service marks owned by MIPI Alliance. All PCI-SIG specifications are registered trademarks or trademarks of PCI-SIG. All other trademarks are the property of their respective owners.

