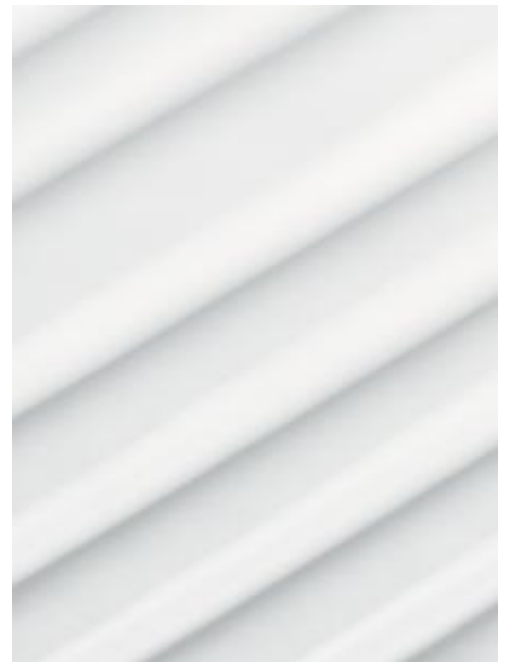


Chip News

August 2026

반도체 산업 최신 이슈를
카드뉴스로 한눈에 확인하세요.



반도체 설계 AI 모델 경쟁, KIMI K3 vs Nemotron 3 Ultra

Chip Design AI Model, KIMI K3 vs NVIDIA Nemotron 3 Ultra

인텔, 스타트업에 RTL 레벨 CPU IP 라이선싱

Intel Licenses an RTL Level x86 Core IP to the Start up

라피더스, 자체 플랫폼에 케이던스 AI 에이전트 도입

Rapidus Brings Cadence's AI Agent into Its Raads Platform

반도체 설계 AI 모델 경쟁, KIMI K3 vs Nemotron 3 Ultra

K3, 48시간 자율 실행으로 45nm 칩을 설계·검증

Nemotron, ACE-RTL 에이전트로 CVDP 평균 통과율 97.1%

구분	Kimi K3	Nemotron 3 Ultra
지능 지수	57	38
출력 속도	34.2 tokens/s	200.8 tokens/s
100만 토큰 비용	\$2.31	\$0.58
칩 설계 실증	48시간 자율 실행, 4mm ² 칩	ACE-RTL, CVDP 97.1%
강점	장기 자율 설계 완주	RTL 반복 효율 (토큰 71% ↓)

Source : Moonshot AI, NVIDIA

Moonshot AI에서 공개한 Kimi K3는 오픈소스 EDA와 Nangate 45nm 라이브러리로 48시간 자율 실행 만에 4mm² 칩을 설계·검증했습니다. 100MHz 타이밍 클로징, 디코드 8,700 tokens/s입니다.

NVIDIA의 Nemotron 3 Ultra는 ACE-RTL 에이전트와 결합해 CVDP 9개 과제 평균 97.1%, 디버깅 100%를 기록했고 반복당 토큰은 71% 적습니다. 칩 설계 도메인 특화 AI 모델의 경쟁도 점차 심화되고 있습니다.

Chip Design AI Model: KIMI K3 vs Nemotron 3 Ultra

K3: a 4 mm² chip designed and verified in a 48-hour autonomous run
Nemotron + ACE-RTL: 97.1% average pass rate on CVDP

Metric	Kimi K3	Nemotron 3 Ultra
Intelligence Index	57	38
Output speed	34.2 tokens/s	200.8 tokens/s
Cost per 1M tokens	\$2.31	\$0.58
Chip design proof	4 mm ² chip, 48-hour run	ACE-RTL, 97.1% on CVDP
Edge	Long-horizon autonomy	RTL iteration (71% fewer tokens)

Source : Moonshot AI, NVIDIA

Kimi K3 built, optimized, and verified a 4 mm² chip in a single 48-hour autonomous run, using open-source EDA tools on the Nangate 45nm library — closing timing at 100 MHz with 1.46M standard cells and an INT4 MAC array.

Paired with the ACE-RTL agent, NVIDIA's Nemotron 3 Ultra averages a 97.1% pass rate across nine CVDP categories (100% on debugging) while using 71% fewer tokens per iteration than Kimi K2.6.

인텔, 스타트업에 RTL 레벨 CPU IP 라이선싱

기존 ISA 라이선스를 넘어 CPU RTL까지 개방
IDM에서 IP 플랫폼으로, Arm식 생태계 전략일지



Source : Intel

인텔이 지난 5월 설립된 스타트업 로자익랩스에 아톰의 x86 명령어 집합(ISA) 구현 기술과, EDA 도구로 칩을 설계할 수 있는 RTL을 함께 제공했습니다. 외부 기업이 인텔 CPU 코어를 기반으로 자체 프로세서를 설계하도록 허용한 첫 사례입니다. 다만 제공된 RTL의 수준과 설계 변경 허용 범위는 아직 공개되지 않았습니다.

RISC-V 기반의 CPU 기업인 Arm과 SiFive와 같은 좀 더 개방적인 형태의 IP 라이선스를 통해 생태계를 확장하고자 하는 전략으로 보입니다.

Intel Licenses an RTL Level x86 Core IP to a Startup

Beyond ISA licensing — Intel hands over the CPU blueprint (RTL)
From IDM to IP platform, a step toward Arm's ecosystem model



Source : Intel

Intel gave Rosaic, a startup founded this May, Atom's x86 instruction set architecture (ISA) implementation along with the RTL needed to design a chip in EDA tools. It is the first case of an outside company being allowed to design its own processor based on an Intel CPU core. The level of RTL provided and the scope of permitted design changes have not yet been disclosed.

It reads as a strategy to expand the ecosystem through more open IP licensing, in the manner of Arm and RISC-V CPU vendor SiFive.

라피더스, 자체 플랫폼에 케이던스 AI 에이전트 도입

파운드리 설계 플랫폼에 EDA 에이전틱 AI를 직접 심은 첫 사례
초기 설계부터 제조 전 최종 검증까지, 소요 시간 최대 절반 목표



Source : Rapidus, Cadence

라피더스가 케이던스의 이노스택(InnoStack) AI 슈퍼 에이전트를 자사 설계 플랫폼 라즈(Raads)에 통합합니다. 파운드리가 고객용 설계 플랫폼에 EDA 업체의 에이전틱 AI를 직접 심는 첫 사례로 추정됩니다.

이노스택은 배치·배선, 타이밍, 전력 분석 등 설계 후반부 툴을 AI가 스스로 실행해 최적안을 찾습니다. 2나노 기술로 2027년 양산을 노리는 라피더스가 플랫폼을 통한 설계 생산성을 경쟁력으로 내세우고 있습니다.

Rapidus Brings Cadence's AI Agent into Raads

First foundry to embed an EDA vendor's agentic AI in its platform
Targeting up to a 2X cut in design-to-signoff time



Source : Rapidus, Cadence

Rapidus is embedding Cadence's InnoStack AI Super Agent into Raads, its AI-agentic design solution — likely the first foundry to put an EDA vendor's agentic AI directly in a customer design flow.

InnoStack autonomously drives back-end tools to converge on an optimal design. For a latecomer aiming at 2nm production in 2027, design productivity is the customer pitch.