

ISART 2026 · MODERATED PANEL

Evaluation Frameworks for ML vs. Classical Radio Solutions



How should engineers decide, validate, and deploy ML alongside classical methods in real-time radios, digital baseband, and sensing systems?

NG

Nada Golmie
NIST

AD

Aleksandar Damnjanovic
Qualcomm

KC

Kaushik Chowdhury
UT Austin

DC

Danijela Cabric
UCLA

Moderator

KB

Kumar Balachandran
Ericsson

When do we use ML — and when classical still wins

Start from the property of the problem, not the novelty of the model.

SUGGESTED QUESTIONS

- Q₁** Which case justifies ML: trusted data with no good analytical model, an intractable optimization, or a genuinely complex new task? ↪ *Which of these is present before training begins?*

- Q₂** Where should classical approaches remain the default—matched models, tight PHY loops, massive-MIMO IQ processing, or a strong closed-form estimator? ↪ *Name the boundary where that answer changes.*

- Q₃** Is the right architecture ML, classical, or hybrid—for example, a learned block inside a physics-based pipeline?

Audience Poll

Should we try to use ML? Is there enough uncertainty
in most modeling to turn into additional
performance?

Evaluation, certification, and fair comparison

A black-box result is not deployable evidence until it survives the lab, field, and full system.

SUGGESTED QUESTIONS

- Q 1** Without analytical bounds, how should a black-box ML radio be tested in the lab and monitored after deployment? ↳ What triggers recovery or fallback in the field?
- Q 2** What must be standardized: public datasets and metadata, a testing harness, system-level metrics—or all three?
- Q 3** How do we prevent benchmark overfitting while preserving reproducible comparisons and vendor differentiation? ↳ Where should standards stop and implementation begin?

Edge latency, model placement, and hardware

Training and inference live on different clocks—and on different parts of the network.

SUGGESTED QUESTIONS

- Q 1** Which functions belong on device, base station, or higher in the RAN—and what latency budget decides placement?
- Q 2** When data movement, buffering, or inference arrives one millisecond late, is compression enough—or should the system fall back to classical? *↳ What is the recovery path when the model is stale?*
- Q 3** Do purpose-built accelerators make the case, or do model size, energy, and retraining cadence still defeat deployment economics?
- Q 4** Are sensing and ISAC better candidates because tens-to-hundreds of milliseconds offer more slack than tight communications loops?

Digital twins: Real-time situational awareness

A useful twin is bounded by the decision it supports and the feedback that keeps it current.

SUGGESTED QUESTIONS

- Q 1** How do we close the sim-to-real gap: field measurements, continuous feedback, online calibration, or explicit uncertainty bounds? ↳ How often must the twin be refreshed?

- Q 2** When the physical world changes faster than the twin, should the system adapt, forget, or fall back to classical?

Audience Poll

Are digital twins of a system necessary for advancing
the state of the art?

Where would you deploy ML today?

Two-minute verdict: ML, classical, or hybrid—plus the deciding metric and the condition that breaks the choice.

RF fingerprinting

DC

Owner: Cabric · UCLA

Neural receiver / channel estimation

KC

Owner: Chowdhury · UT Austin

ISAC and location

ND

Owner: Damnjanovic · Qualcomm

DPD and beam management

AD

Owner: Damnjanovic + Cabric

If time is tight, take only the first two. Do not sacrifice the protected 15-minute audience window.

Audience questions

What do we carry out of the room?