

(As Prepared for Delivery)

Reflections on Three Eras of Wireless Technology: Efficiency, Structural Margins, and Network Resilience

Keynote Address by Dale N. Hatfield

Before the International Symposium on Advanced Radio Technologies (ISART 2026)

Boulder, Colorado

August 13, 2026

Introduction

Thank you, David and Jeremy, for the introduction, and thank you both for your work in reviving the ISART conference after its recent hiatus.

Standing here in Boulder this afternoon brings two thoughts to mind. First, I am gratified to see so many of my friends and colleagues in the spectrum management field in the audience, especially these days, when it is difficult for me to travel. Second, it reminds me of how important this forum remains. For decades, ISART has served as a vital space where our community can come together to discuss and advance the state of radio technology.

The issues we discuss here are not merely technical or academic. Failures in the systems that depend on reliable access to and use of spectrum can trigger cascading economic and social consequences and directly threaten public safety, critical infrastructure, and national security. That is why I am convinced that this collaborative gathering is more necessary than ever.

I first spoke at ISART in September 2000, while serving as Chief of the FCC's Office of Engineering and Technology. I returned in 2002 to discuss "Spectrum Management in the New Millennium," and again in 2004 to examine the challenges associated with relying more heavily on marketplace forces in spectrum management. When David and Jeremy asked me to speak, they suggested that I look back over the trajectory of spectrum management and technology, from my early days in the field through my previous presentations at this conference. I agreed. As part of that process, I reviewed my earlier ISART presentations. That review led me to think about this history in terms of three distinct eras in wireless networks.

The first was an era dominated by comparatively rigid physical limitations on such networks. The second has been an era of extraordinary digital and software-enabled agility. The third, the era we are entering now, raises some pressing architectural questions that I have been grappling with over the past year in my own research.

The question connecting those three eras is straightforward: as we extract ever greater efficiency from spectrum, are we also removing or obscuring the very margins, such as physical guard bands, spatial buffers, and engineering trade space, on which network resilience depends?

The Era of Hardwired Physical Limitations

To understand our current state of hardware and software agility, it helps to look back at the world I stepped into when I joined the Central Radio Propagation Laboratory (CRPL), a predecessor of ITS, here in Boulder in 1963. That was more than sixty years ago.

As a young amateur radio operator, experimenter, and engineering student in that hardware-dominated era, I had encountered many technical and operational limitations directly. If I wanted to change frequency, I had to pull one piezoelectric crystal in its holder and replace it with another. If I wanted to install a vacuum-tube two-way radio in a vehicle, the equipment could take up much of the trunk. The transmitter drew so much current that, when I keyed the microphone at a stoplight, the engine might labor and the headlights might dim.

The condition of the equipment was often equally apparent. I could look at the plates of a vacuum-tube amplifier and, if the plates were glowing red, I knew it was being driven too hard. I could tune across the channels and hear both co-channel and adjacent-channel interference. And, in some cases, my sense of smell told me that I had burned out a component before any instrument did.

Now those physical limitations had direct consequences for how we designed and operated wireless systems. The equipment was bulky and power-hungry. Frequency stability was limited. Filters could be physically large, channelization was coarse, and maintenance was a continuing burden.

We accommodated those limitations by building generous margins into wireless systems. Imperfect frequency stability and transmitter spectral purity justified wider channel spacing. Limited receiver selectivity justified greater frequency or geographic separation. Conservative link budgets provided margin for fading, component variation, and propagation conditions we could not predict reliably.

Those cushions could be inefficient, but they were tangible. An engineer could point to the filter, the guard band, the geographic separation, or the excess link margin. Degradation was often audible as noise, distortion, fading, or reduced intelligibility. So we frequently had both physical cushions protecting the system and physical indications that told us when the system was approaching the limits of reliable operation.

Our Current State of Software Agility

When I returned to the FCC in 1997, after time in the private sector and academia, an experienced Commission engineer who was preparing to retire approached me as I was entering FCC headquarters. He offered me a piece of advice. He told me that, if I wanted to understand where radio technology was going, I should pay close attention to software-defined radio.

He was right.

Functions that had once been fixed in crystals, filters, oscillators, modulators, and other hardware were increasingly being moved into software. A radio once designed for one frequency range, one waveform, or one operating mode could be reconfigured, sometimes almost instantaneously, by changing its software.

At the same time, in the late 1990s, many of us were deeply concerned about a looming “spectrum drought.” Demand for wireless capacity was increasing rapidly. Broadband services were expanding. New devices and applications were emerging. Yet large portions of the spectrum were locked into allocations, assignments, and regulatory structures designed for an earlier technological era.

In one of my earlier ISART presentations, I discussed marketplace mechanisms, secondary markets, unlicensed access, spectrum sharing, software-defined radio, and more flexible rights of use. All were subjects of active discussion throughout the spectrum-management community at the time. The goal was not simply to make additional spectrum available through new allocations or assignments. It was to obtain more useful communications capacity by using the time, frequency, and space dimensions more efficiently.

The demand for greater communications capacity from the spectrum resource arose among nearly every class of user, especially the rapidly growing commercial wireless sector. The engineering community responded more successfully than many of us could have imagined. Wideband digital sampling, high-speed signal processing, advanced coding, adaptive modulation, multiple-input multiple-output systems, interference cancellation, dynamic power control, geolocation, databases, sensing, and software-defined waveforms have transformed what can be accomplished in a crowded electromagnetic environment.

We can reuse frequencies across small geographic areas. We can coordinate sharing among systems with widely different characteristics. We can change coding, modulation, power, bandwidth, and routing as conditions change. We can support data rates and device densities that would once have seemed impossible.

These advances have not merely increased efficiency. Many have also improved network resilience. A radio that can change frequency, alter its waveform, reroute traffic, use diversity, or adapt its coding can be much more capable of surviving interference and changing conditions than a rigid radio designed for one channel and one operating mode.

The issue, therefore, is not that software agility is inherently fragile, nor that we should return to inefficient analog systems. Rather, the challenge as we look toward AI-native architectures is ensuring that our growing software sophistication does not blind us to the physical limits underneath.

The Perils Ahead and the Role of ITS

This brings me to the more subtle issue at the center of my remarks. Increased efficiency can come at the price of resilience. Protective margins may be removed or hidden, or replaced by complex mechanisms whose own dependencies are not fully understood. In short, the margins protecting modern systems have changed form.

Some of the margin once supplied by relatively coarse frequency, spatial, and equipment separation has been shifted into coding, adaptive control, coordination databases, sensing systems, interference cancellation, receiver processing, network orchestration, and software more generally.

I think it is also useful to distinguish between passive physical margin and active, complexity-dependent margin. A guard band consumes spectrum. It does not require a software update, a timing source, a database connection, a sensor, or agreement among several networks. An adaptive interference-management system can use the same spectrum much more efficiently. The protection it supplies depends on all of those things and on their interactions.

Clearly, active margins can produce a substantial net benefit. They also introduce dependencies and failure surfaces of their own. Diversity provides resilience only when the supposedly diverse paths are genuinely independent. Two links on different frequencies may still share the same tower, power, timing, backhaul, software, or control plane. Automation can respond faster than people. It can also spread the effects of an erroneous response more quickly and over a much larger area.

The question is not whether such complexity is good or bad. It is whether the reduction in the original risk is greater than the additional risk introduced by the mechanism intended to control it. Answering that question calls for system-level analysis of interfaces, dependencies, transition states, and common-mode failures, rather than merely demonstrating that each individual feature works as designed. Simply put, we cannot evaluate the safety of a complex system by testing its software components in isolation; we must measure how the entire architecture behaves under stress.

There is another question: can we see the margin being consumed? Software agility changes what the operator and user can see. In an older analog system, a degrading signal often announced itself through increasing noise, fading, or distortion. On a modern digital link, error

correction, retransmission, adaptive modulation, buffering, diversity, and power control may keep the application working while the physical environment deteriorates underneath it.

Signs of strain may still be present: declining signal-to-interference-plus-noise ratio, rising retransmissions, more robust but less efficient coding, stressed synchronization, or a receiver approaching blocking, desensitization, or a nonlinear threshold. Those same adaptations may keep the warnings from the user, and sometimes from the people responsible for the health of the network.

A fuller picture of network resilience therefore has two parts. One concerns service: availability, throughput, latency, packet delivery, and whether the application still works. The other concerns the physical foundation: signal quality, interference, noise, synchronization, receiver stress, retransmission burden, and coding margin. If we keep only the first view, we may mistake continued service for a healthy system.

Observing that physical foundation is essential, but so is the ability of the receiver to withstand the conditions it encounters. Receiver performance has therefore been a concern of mine for most of my professional life. I learned that lesson as a teenage amateur radio operator more than seventy years ago, when my ham radio transmissions interfered with a neighbor's reception of television Channel 2. The solution required attention not just to the transmitter end, but to the receiving end as well. The result was a high-pass filter at the neighbor's television receiver and a substantial low-pass filter at my transmitter, where the filter had to handle much higher power.

That small episode taught me an enduring lesson: interference is a system-level phenomenon involving the transmitter, the receiver, and the environment in which they operate. Yet our regulatory system has concentrated much more heavily on controlling emissions than on ensuring receiver performance. A transmitter may comply fully with its authorization and service rules and still overwhelm a receiver that lacks adequate selectivity, dynamic range, or resistance to blocking and desensitization.

This asymmetry bears directly on the relationship between efficiency and resilience. As we pack systems more closely together, reduce guard bands, and rely more heavily on coordination and signal processing, we use up margins that once helped protect receivers from both expected and unexpected conditions. We gain spectrum efficiency. Without comparable attention to receiver performance, monitoring, and enforcement, we may optimize the system beyond its ability to absorb a failure.

Allocation, service rules, equipment authorization, receiver performance, monitoring, and enforcement work best together as complementary parts of a larger system. Effective interference response is also a process rather than a single act. It proceeds from detection through investigation, analysis, mitigation, remediation, documentation, and feedback. Otherwise, we may correct the immediate problem without learning enough to prevent its recurrence.

Nor will the consequences always remain local. A receiver failure at one cell site or control point can spread through shared timing, backhaul, software, databases, and automated coordination. What begins as a local physical-layer problem may become a regional or even national service disruption. This brings me back to our hosts. ITS and its predecessors have made an important national contribution by testing software assumptions against physical evidence. Their long tradition of propagation research, spectrum measurement, radio-noise studies, equipment characterization, and model validation is increasingly important.

One area in which independent physical evidence is especially important is the man-made radio-noise environment. That environment has changed and, in important respects, changed for the worse. During the hardware era I described earlier, automotive ignition noise was among the most familiar sources, whether I encountered it in amateur radio, two-way mobile systems, or my work at CRPL. Improvements in ignition systems have reduced that particular problem. At the same time, the source environment has changed dramatically. For example, switched-mode power supplies are now ubiquitous. They offer high efficiency, small size, and light weight in electronic devices of all types. They also produce noise with different spectral, temporal, and geographic characteristics.

Clearly, what is missing is not measurement expertise; much of that expertise is here in this room. What would be valuable, in my opinion, is sustained support for rigorous, long-term measurements that can establish reliable baselines, reveal trends, identify changing sources, and help us direct scarce technical and enforcement resources where they will do the most good. That is precisely the kind of work for which ITS has been, and still is, exceptionally well suited.

The same principle applies to receiver behavior, transmitter imperfections, active and passive intermodulation, propagation models, and automated spectrum-sharing systems. As spectrum management becomes more dependent on databases, models, and machine-to-machine decisions, an independent capability to compare those abstractions with the electromagnetic environment becomes more—not less—important.

Looking ahead, the growing use of physics-informed machine learning is one response to this need to keep our models connected to physical reality. I have been particularly interested in work on Hamiltonian and structure-preserving approaches. In plain terms, a Hamiltonian represents the energy and dynamics of a physical system. The associated symplectic structure describes relationships in the system's evolution that should be preserved over time.

The attraction is not that these methods eliminate problems of propagation or spectrum management, but that they may help us understand and manage those problems more effectively. They seek to guide learning by incorporating known physical laws and relationships, rather than relying entirely on statistical fit. Where applicable, that may improve long-term stability and make the model less likely to drift away from known physical behavior.

Nevertheless, the same complexity test still applies. Every model embodies assumptions. A physics-informed model may preserve one relationship while omitting another, or faithfully

enforce an incomplete representation of the physical environment. Measurement and independent validation therefore remain essential.

Conclusion

In conclusion, and well within my lifetime, we have traveled a remarkable distance from the world of crystals, vacuum tubes, large filters, and dimming headlights. Software agility, adaptive radios, advanced coding, and automated coordination have allowed us to obtain vastly more communications capacity from the spectrum. We should be proud of that achievement and continue to build upon it.

We should also recognize exactly where the margins protecting our systems now reside. Efficiency and resilience are not opposites, but neither are they automatically aligned. Resilience sometimes requires slack, independence, physical observability, and room for recovery. What looks like unused capacity may be performing a protective function. What looks like successful software adaptation may be concealing physical deterioration.

If there is one thought I hope to leave with you this afternoon, it is this: excessive optimization is the enemy of network resilience.

Thank you very much for your attention.