

Before the Hardware Exists

What reliability engineers should do instead of spending early program time on prediction models

Kirk A. Gray | Accelerated Reliability Solutions, L.L.C.

“Prediction is very difficult, especially if it is about the future.”

— attributed to Niels Bohr

A common question comes up whenever I challenge the value of spending early engineering time on reliability prediction models: well, if we are not going to calculate a predicted failure rate before the product exists, what should we do instead?

My answer is simple: use the time to learn where the real risk is. Long before a complete hardware system is available for HALT or system-level stress testing, there is plenty of useful reliability work to do. In fact, this early work often determines whether the later test program finds meaningful design and manufacturing weaknesses or merely produces a report.

Prediction models may look precise, but precision is not the same as truth. If the model does not represent the actual causes of field failures, it can consume time without reducing risk. The better approach is to build evidence: study past failures, understand margins, compare supplier strength, plan useful monitoring, and prepare the system so weaknesses can be found before customers discover them.

1. Start with the evidence from earlier products

The most valuable reliability information is usually not in a prediction handbook. It is in the history of your own products, your suppliers, and your customers. What failed? Why did it fail? Was it a weak design margin, a manufacturing variation, a supplier change, a connector issue, a cracked solder joint, a marginal power rail, a thermal corner, a timing issue, or a misuse condition that should have been anticipated?

A predecessor-product review should not be a blame exercise. It should be a risk-reduction exercise. The purpose is to make sure known mechanisms have been removed, improved, or deliberately monitored in the new design. If a previous product had field failures that were never fully understood, that uncertainty belongs on the reliability plan for the next product.

- Review field returns, warranty data, failure analysis reports, customer complaints, and no-fault-found cases.
- Separate true wear-out mechanisms from design-margin problems, process escapes, supplier variation, and application abuse.
- Turn each known mechanism into a design question, a test condition, or a monitoring requirement for the new product.

2. Build the reliability test strategy before the hardware arrives

A useful reliability test plan is not just a calendar of tests. It is a plan for exposing the most likely weaknesses in the most useful order. Early in the program, the team should decide which assemblies matter most, which stresses are most likely to expose low margins, what functionality must be operating during stress, and what failure signatures need to be captured.

This is also the right time to connect the reliability plan to the FMEA. The FMEA should help identify critical functions, likely failure effects, and areas where a small design or process shift could create a large field

consequence. The test plan should then focus on those areas instead of treating every subsystem as equally risky.

- Define the most important functions that must be active during stress.
- Decide what will be tested at the component, board, subsystem, and system levels.
- Identify which stresses matter: temperature, thermal cycling, vibration, voltage margining, load variation, clock speed, bus frequency, humidity, or combinations of these.
- Plan the order of testing from least destructive to most informative, so the same units can teach you as much as possible.

3. Plan the monitoring as carefully as the stress

In HALT and other accelerated evaluations, the response of the unit under test is more important than the chamber set point. A test that applies stress but does not monitor the right signals can miss the failure mechanism entirely. Worse, it can create a false sense of confidence.

Before the final system exists, engineering teams can still plan how the system will be instrumented. Thermocouples, multichannel data loggers, accelerometers, spectrum analysis, electronic loads, power monitors, diagnostic software, event logs, and high-speed captures may all be needed. For power supplies, programmable electronic loads are especially useful because voltage regulation under stress and load variation can reveal weakness long before a simple pass/fail test does.

Subsystem testing may also require practical preparation: extension cables, breakout fixtures, remote sensors, external loads, software diagnostics, remote boot methods, or temporary ways to operate one assembly while another remains outside the chamber. These details sound mundane, but they often determine whether the test produces useful engineering information.

4. Test supplier subsystems before the full system comes together

Large electronic systems are built from many assemblies that do not all arrive at the same time. There can be a long gap between receiving a power supply, motherboard, graphics card, storage assembly, interface board, or display module and having the complete system ready for final validation.

That gap is not dead time. It is an opportunity. Vendor-supplied subsystems can be tested and compared empirically. Instead of assuming that two suppliers are equivalent because they meet the same data-sheet specification, we can compare how much margin each design actually has.

- Measure upper and lower empirical operating temperature limits.
- Measure vibration operating limits or the highest safe stress completed.
- Check voltage regulation under variable loads and temperature extremes.
- Evaluate whether failures are recoverable, destructive, intermittent, or tied to a specific functional mode.
- Compare multiple samples when possible so the team can see variation, not just one lucky or unlucky unit.

Example: using empirical limits to compare power-supply suppliers

The following simplified example shows why empirical limits can be more useful than a paper prediction. If three power supplies all meet the purchasing specification, the supplier with the largest demonstrated margin may still be the better reliability choice, assuming price, production capability, and supplier quality are otherwise comparable.

Measured item	Supplier A	Supplier B	Supplier C
Upper operating limit	120 °C	100 °C	115 °C

Lower operating limit	-50 °C	-40 °C	-40 °C
Thermal operating range	170 °C	140 °C	155 °C
Relative margin rating	Best	Good	Better

Table 1. Example comparison of empirical thermal operating margins for three power-supply suppliers.

The table does not prove that Supplier A will always be more reliable. It does show that Supplier A has more demonstrated thermal operating margin in this example. That is practical evidence the team can use during supplier selection, especially when it is combined with failure-mode observations, cost, manufacturing capability, and field history.

5. Create margin maps, not just pass/fail reports

A pass/fail test answers one narrow question: did this unit survive this condition on this day? A margin map answers a more useful question: where does the product begin to lose function as stresses are varied?

For modern electronics, the useful boundary may be multidimensional. Temperature, input voltage, load, clock rate, bus frequency, vibration, and software activity can interact. A board that operates at one voltage and temperature may become intermittent when the voltage is reduced, the temperature is raised, and the bus frequency is increased. That is exactly the kind of margin problem that may later look like a mysterious software fault in the field.

Artificial intelligence and modern data tools can help here, not by pretending to predict a universal failure rate, but by helping engineers visualize multi-stress boundaries and detect shifts over time. The goal is to build a practical map of the product's operating strength, then use that map as a benchmark.

A two-dimensional view of measured operating limits

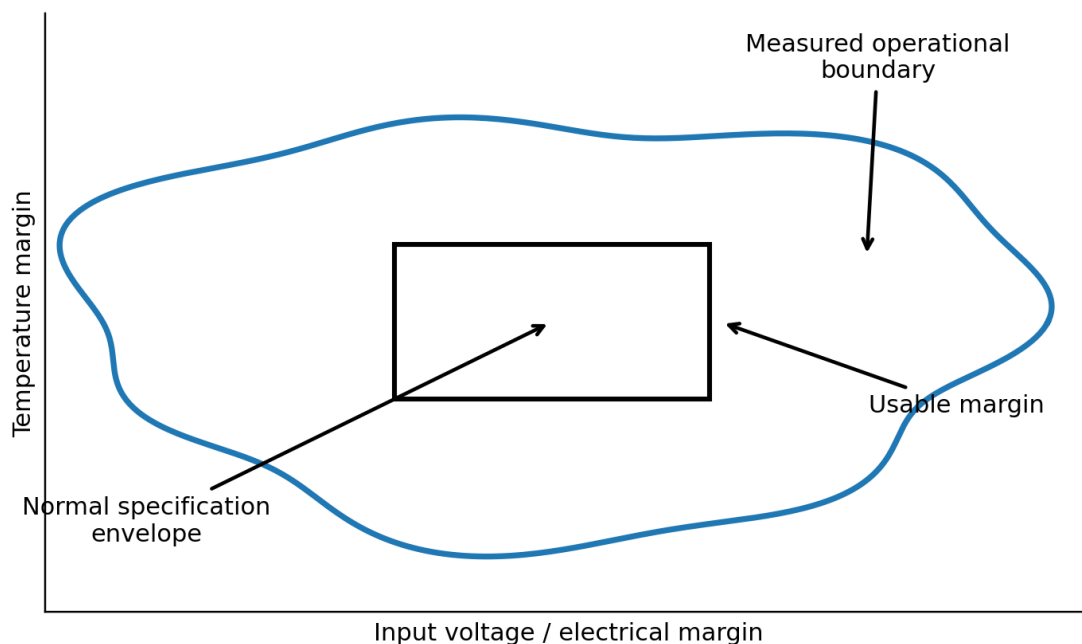


Figure 1. Conceptual two-dimensional stress map showing a measured operating boundary outside the normal specification envelope.

Tracking margins over time can reveal production drift

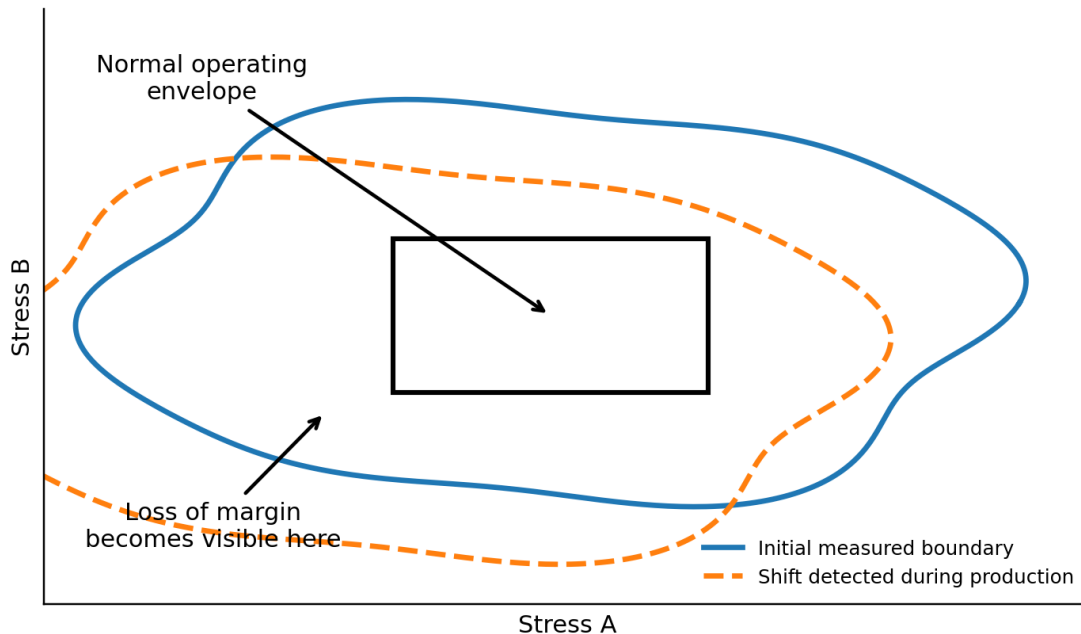


Figure 2. A production shift in the measured boundary can reveal loss of margin before it becomes a field failure pattern.

6. Turn early learning into production controls

The work done before full system test should not disappear when the design is released. The same empirical limits and failure signatures can become the basis for production monitoring, HASS, HASA, or ongoing reliability audits.

For example, if voltage regulation under thermal stress is a sensitive discriminator for a power supply, that measurement may become a useful production indicator. If a processor board becomes marginal only near a combination of high temperature, low voltage, and bus frequency, that combination may become a valuable audit condition. The goal is not to abuse good products. The goal is to detect a meaningful loss of strength in time to prevent defective populations from reaching customers.

- Use HALT results to define safe HASS or HASA stresses after proof-of-screen work is complete.
- Monitor the parameters that actually moved near the operating boundary, not just the easy measurements.
- Use supplier changes, engineering changes, process changes, and component substitutions as triggers to recheck margins.
- Treat shifts in limits as early warnings, even when units still pass normal functional tests.

7. A practical pre-hardware reliability checklist

Task	Question to answer
Predecessor review	What failed before, and have those mechanisms been removed or controlled?
FMEA connection	Which functions and failure effects deserve the most attention?
Subsystem plan	Which assemblies can be tested before the complete system is ready?
Stress plan	Which stresses are most likely to reveal low margins?
Monitoring plan	What signals must be measured during stress to detect intermittent or marginal behavior?

Vendor comparison	Do competing suppliers have the same demonstrated strength, or only the same specification?
Margin map	Where is the actual operational boundary under combined stresses?
Production control	How will we know later if the product or supplier process has become weaker?

Table 2. A simple reliability work plan before complete system hardware is available.

Conclusion: do useful work before the complete system exists

There is no need to wait passively for a complete hardware system before doing meaningful reliability engineering. And there is little value in spending scarce early program time on prediction models that do not reflect the real causes of most electronics unreliability.

A better use of time is to build evidence. Learn from past designs. Understand supplier and subsystem margins. Plan the monitoring. Create stress-limit maps. Identify the measurements that reveal weakness. Then use that knowledge to guide HALT, HASS, HASA, and production audits.

The purpose of reliability engineering is not to create a number that looks authoritative in a report. The purpose is to prevent customers from discovering weaknesses that engineering could have found earlier. Before the complete system exists, the best reliability work is not prediction. It is preparation, measurement, comparison, and margin improvement.