

There is a particular moment that shows up in almost every access control project. A door that looked fine on paper becomes political in the field. Someone asks a question that seems simple until you realize it changes the entire design: "If the power fails, what do you want this door to do?"

That question is really about philosophy, risk tolerance, and building operations. It is also where people get tripped up by the terms fail-safe and fail-secure. Those labels sound like they map cleanly to "good" and "bad", but in practice the right choice depends on life safety goals, operational reality, and the failure modes your site can actually tolerate.

Below is a practical way to decide between fail-safe and fail-secure locks, with the trade-offs spelled out, including the edge cases that cause last-minute redesigns.

Start with what "failure" means on your site

"Power outage" is the most obvious failure, but it is not the only one. When you talk about fail-safe versus fail-secure, you are really talking about what happens when the locking mechanism loses a controlling condition.

That controlling condition might be:

- electrical power
- an access control signal (card reader, credential validation)
- a monitoring circuit
- the controller's ability to command the lock
- a communication link between the controller and the system head-end

You do not have to predict every failure, but you do need to decide what you are optimizing for. A hospital corridor under fire code constraints is optimizing for evacuation and smoke movement. A secure server room is optimizing for theft resistance and containment. A warehouse with a lot of foot traffic is optimizing for flow and reducing the chance that a random incident traps someone in a dead-end.

If you approach the decision as "what should happen when something goes wrong," you can make the terminology serve the real-world goal, instead of the other way around.

The core behavior: fail-safe versus fail-secure

Most of the confusion comes from how the industry phrases these terms.

- **Fail-secure** locks are designed to stay locked when power or control is lost. In other words, the default state under failure is "deny entry."
- **Fail-safe** locks are designed to unlock when power or control is lost. The default state under failure is "allow egress," which usually means the door becomes operable for people to get out.

In a perfect model world, fail-safe supports egress during an outage, and fail-secure supports security during outages. In the real world, what matters is which risk you are willing to accept, and whether your door control approach still supports safe movement and required unlocking during emergencies.

One practical note that I learned the hard way: teams sometimes treat "fail-safe means unlock" as a blanket statement and then wire the alarm and release logic inconsistently. If the system can unlock the door via multiple

paths (fire alarm, emergency release, manual egress hardware), you need to confirm that the actual event path lines up with the building's life safety strategy.

Decide based on the door's job, not the hardware label

The phrase "door's job" sounds obvious, but it changes your choices when you examine the intent behind the opening.

Ask what the door is primarily controlling:

- **Egress and emergency travel:** doors in corridors meant for evacuation, stair access, and critical egress paths.
- **Normal access to restricted areas:** offices, labs, or floors where people can be prevented from entering without creating an evacuation hazard.
- **Perimeter or asset protection:** doors protecting high-value areas, secure storage, data rooms, or areas where unauthorized entry is a major concern.
- **Segregation and operational control:** doors used to manage traffic patterns, separate hazards, or enforce process separation.

When a door is part of a required means of egress, the design team is typically optimizing for people leaving safely, even if it means the lock releases during failure conditions. When a door is part of a restricted security boundary, the organization often prioritizes keeping unauthorized people out, even if it means the lock remains engaged when power fails.

But there is a third variable people forget: you are rarely choosing between only "unlocked" and "locked." You are choosing between multiple behaviors across multiple conditions, like alarm release, emergency egress, and scheduled access.

That is where the real decision becomes more nuanced.

Life safety tends to drive fail-safe choices, but confirm the whole emergency sequence

In many building types, life safety requirements strongly influence lock behavior. During fire or life safety events, doors often need to unlock, release, or allow free egress. In that scenario, fail-safe locks can simplify the story: when control power is lost, the door defaults toward allowing people to exit.

However, this does not mean fail-safe is always correct for every life safety opening. Sometimes doors need to remain controlled for compartmentation, smoke management, or fire-rated behavior, and the hardware selection needs to support the door's fire strategy.

What I've seen work reliably is not just choosing the lock type, but making sure the entire emergency sequence is coherent:

- If the fire alarm activates, does the door release as required?
- If power fails during an alarm event, does the release still occur?
- If the system controller is down, do local release devices still operate correctly?
- Are there any conditions where the door could remain locked even though it should be open for egress?

Even if your instinct says "fail-safe," the system might still need an explicit emergency release path. Conversely, even if you choose fail-secure for security reasons, you still have to ensure that emergency egress requirements

override normal access control. That override is usually handled by fire alarm interfaces and egress hardware, not by assuming the lock logic will magically match code intent.

If you are working with an AHJ (authority having jurisdiction), it is worth validating early. Lock logic details are exactly the kind of thing inspectors and fire marshals want to see mapped clearly.

Security and containment often favor fail-secure, but watch the evacuation path

For restricted areas that are primarily about preventing unauthorized entry, fail-secure defaults can be attractive. When power fails, the door stays locked, which reduces the “open door during outage” window that attackers and opportunists sometimes look for.

This can be a good approach for:

- server rooms and network closets
- labs with controlled access and sensitive equipment
- vaults and secure storage
- areas with managed visitors, where letting everyone in during an outage would undermine policy

But your evacuation path still matters. If a door is on an egress route, keeping it locked during an outage can become an operational hazard even if the lock itself is designed for security.

The fix is usually not “switch to fail-safe everywhere.” The fix is to align:

1. What the door is allowed to do during normal conditions,
2. How emergency egress is supported,
3. What happens during power and controller failures.

In real deployments, fail-secure doors often require careful integration with:

- egress hardware that provides a guaranteed path out
- emergency release circuits that override locking during alarm events
- local manual hardware that can operate even if the system is partially down
- monitoring logic so failures and forced egress are visible and actionable

If you pick fail-secure for a security door but do not confirm that people can always get out, you end up with the worst kind of compliance risk: a door that is technically “secure” but can trap occupants during the exact kind of failure that should be survivable.

The human factors piece: what people will do during an outage

Hardware logic matters, but human behavior during stress is equally important. When people are in a hurry, they tend to treat doors as binary objects: push, pull, try again, and look for someone who can help.

During a power outage, a fail-secure door that stays locked can cause confusion and delays. In some facilities, it is common for staff to have a local procedure, like calling a security desk or using a manual override. That works when trained staff are present and when the procedure is well [commercial access control companies](#) communicated.

During a short outage at a staffed site, people might not even notice because your emergency plan keeps egress clear. During a longer outage at an unstaffed site, a fail-secure default can create bottlenecks, especially in high-traffic corridors and stair approaches.

I remember a case where a facility installed fail-secure locks on doors that were not truly “exit doors,” but were used like shortcuts. On a Saturday outage, the doors stayed locked, and people started pushing harder and waiting. The building was safe, but it created a situation that security and operations were spending the rest of the day managing. The fix was not changing everything to fail-safe, it was correcting the access plan, updating signage, and ensuring the emergency behavior sequence was clear.

So, include operations in your decision. Ask what your staff can realistically do during outages, and how long it takes them to respond.

Operational continuity and maintenance realities

Fail-safe and fail-secure choices are not only about failure states. They also affect day-to-day maintenance.

Locks, power supplies, and controller interfaces all need periodic testing. If your design relies on a specific release behavior during emergency conditions, you will end up testing it. That means your chosen strategy should be testable without turning the building into a fire drill.

There are also power-related edge cases:

- If you use power failover or UPS, the lock may behave differently than expected during the early seconds of an outage.
- Some installations have “brownout” conditions where voltage sag causes intermittent behavior. That can be more annoying than a full outage.
- If you have distributed controllers or local fail logic, you need to know which component actually decides the lock state during failure.

A lot of teams focus on the lock definition and forget the surrounding architecture. The question to keep returning is: during a realistic failure scenario, which component enforces the lock state?

That is the component you need to understand, document, and validate.

A decision framework that works in the field

A clean selection process usually looks less like “pick fail-safe because it sounds safer” and more like a structured risk decision.

One workable approach is to evaluate each door on three dimensions:

1. Egress and life safety impact

How likely is it that someone would need to exit through this opening under stress or during a failure?

2. Security boundary impact

What is the consequence if unauthorized entry is possible during an outage?

3. Override and fallback behavior

Even if the lock defaults one way, do you have guaranteed override paths for emergencies and guaranteed exit mechanisms?

You rarely answer these questions with perfect certainty, but you can reach a defensible selection.

Here is the practical shortcut I use: if the door must always allow people out during the scenarios your building is designed to survive, your system must guarantee that regardless of the lock type label. If it must deny access for containment and the building still provides a reliable exit route, then fail-secure can make sense, provided emergency logic and hardware are properly integrated.

When “fail-safe” and “fail-secure” get mixed in one project

Modern access control systems can be configured so different conditions produce different lock states. You might have a door that is generally secure but unlocks on fire alarm activation, while still remaining locked on loss of normal power. This is where projects get messy if the design documents do not clearly state which event triggers which behavior.

Common mixed scenarios include:

- Normal condition locked, fire alarm releases, power outage keeps locked unless the fire panel triggers local release.
- Normal condition unlocked for scheduled hours, locked outside schedules, but emergency egress always overrides.
- Credential reader present for access control, but mechanical override and egress hardware provide an exit independent of the controller.

In these cases, the distinction between fail-safe and fail-secure becomes less about the lock’s label and more about what your emergency interface and local hardware actually do.

If you are managing a multi-door rollout, treat each door like a small system. Document the exact triggers and outcomes for each door, and avoid assuming that “the system will handle it.”

The checklist I wish more designers used before wiring decisions

This is not a substitute for code compliance or manufacturer guidance, but it prevents many preventable mistakes. Use it when you are about to finalize wiring drawings, interface points, and programming logic.

- Identify whether the opening is part of a required means of egress and confirm the intended emergency behavior with the appropriate stakeholders.
- Define the specific failure scenarios you are modeling: full power loss, controller failure, communication loss, and fire alarm activation.
- Confirm what component controls the lock state during each failure scenario, including any local release hardware.
- Verify that emergency egress is possible even if the lock defaults to locked (for fail-secure) and even if access control power is unavailable.
- Plan how you will test the behavior without disrupting operations more than necessary.

That checklist alone will not make your decision for you, but it forces clarity where teams often rely on assumptions.

Concrete examples to anchor the trade-offs

Example 1: Office floors with controlled doors

Imagine an office building where suite doors need controlled access, but corridors and stairwells are the real egress routes. Many suite doors are security boundaries, and the owner does not want doors opening during routine outages.

A typical outcome: you can choose fail-secure for the suite door lock logic, because egress is not primarily dependent on that door. You then ensure that emergency egress paths exist through required exits and that any emergency release or manual escape mechanism for that specific opening meets the applicable requirements.

The main trade-off is operational confusion during outages. People may hit a locked suite door and assume it is a malfunction. That can be mitigated with signage, a procedure for staff, and system monitoring.

Example 2: A corridor door that people use like an exit

Consider a door in a healthcare or education setting that is technically not the primary exit but becomes the practical exit route during everyday operations. People use it because it is closer.

If you choose fail-secure for security reasons and the door stays locked during an outage, you create a mismatch between real human behavior and intended design. Even if code compliance is met, you may see crowding, frustration, and delayed evacuation movement.

In that kind of environment, fail-safe default behavior or strong emergency override logic tends to reduce friction, because the building's layout makes people treat the opening like an exit.

Example 3: Secure data closet with guaranteed emergency egress override

Now picture a small data closet protected for asset protection. Unauthorized entry is a serious issue. You choose fail-secure so the door remains locked during power loss.

But the closet door still needs to allow safe exit for occupants who are inside. You verify a local exit hardware solution that allows egress even if the lock is in secure mode. Then you integrate the fire alarm release so the door behaves properly during alarm conditions.

This example highlights the key point: "fail-secure" does not mean "dangerous." It means you must engineer the overrides so that emergency egress is not dependent on the access control system remaining powered.

Common edge cases that change the decision

There are a few situations where the usual "fail-safe for egress, fail-secure for security" rule of thumb breaks down or requires extra care.

Edge case: Doors with delayed release expectations

Some facilities want doors to remain locked briefly during certain transitions, then unlock under emergency conditions. If you implement timing logic incorrectly, you might cause the door to remain locked longer than intended.

This is especially risky for doors adjacent to evacuation routes, where even a short delay can become a barrier under stress.

Edge case: UPS and generator behavior

If your lock strategy depends on power loss being immediate, but you provide UPS for controllers or readers, the observed behavior during “outage” may not match the design assumptions.

A door might remain locked longer because the controller stays alive, then suddenly change state when UPS runs down. If your team expects an immediate unlock for safety, you need to confirm how long “power loss” really lasts for the lock logic.

Edge case: Maintenance-induced failures

The failure mode you care about is not only “an attacker cuts power.” It is also “someone miswired a relay,” “a technician replaced a power supply,” or “a door contact failed open.” If your documentation and commissioning tests are weak, a maintenance mistake can turn an intentional fail-safe into fail-secure behavior, or vice versa.

That is why commissioning and testing matter as much as the initial selection.

How to document the decision so the project survives handoffs

Lock decisions tend to fail at handoff. A person picks fail-secure for security reasons, but the fire alarm contractor or installer later wires the release points differently. Or the programming logic changes during integration.

To keep it solid, document three things clearly:

1. **Normal behavior** (who can open it and under what conditions).
2. **Emergency overrides** (fire alarm behavior, local manual egress behavior, and any required release sequences).
3. **Failure behavior** (what happens during controller failure and power loss, not just what happens during a fire alarm).

When these are written in plain language and mapped to the actual wiring and programming points, the decision becomes durable. Teams can test it. Inspectors can review it. Technicians can troubleshoot it.

Practical rule of thumb that stays honest

If you want a simple guiding statement, keep it grounded like this:

- Choose **fail-safe** when your primary goal is ensuring the door defaults toward allowing egress during the kinds of failures you are trying to survive.
- Choose **fail-secure** when your primary goal is denying access during loss of normal control, and you have engineered and verified emergency exit pathways that do not depend on the access control system staying healthy.

That is still not a substitute for code review, door hardware selection, and manufacturer instructions. But it keeps the decision tied to risk, not to terminology.

The final check: can you explain the lock behavior in one minute?

Before you sign off, ask yourself a simple question: can you explain what the door will do when:

- power fails
- the controller fails
- the fire alarm activates

- someone inside needs to exit during stress

If you cannot answer quickly and specifically, the hardware label is not your problem. The system design is not yet clear enough, or the documentation and commissioning plan are missing critical details.

A well-chosen fail-safe or fail-secure strategy does not just meet a requirement. It makes the entire building's behavior predictable, testable, and defensible when something goes wrong.

That predictability is what customers, operators, and inspectors ultimately care about, and it is what prevents the "why did this door do that?" calls long after the ribbon-cutting.