

49 · Testing

Flutter Practice Notes

Contents

49 · Testing

Testing Fundamentals & the Pyramid

Unit Testing, Mocking & TDD

Widget & Golden Tests

Integration & E2E Testing

Testing Integration (Capstone: Test Every Layer + Coverage + CI)

49 · Testing

Introduction

This module covers testing Flutter apps properly: the **testing pyramid** (many fast unit tests, fewer widget tests, few integration/E2E tests) and *why*, **unit testing + mocking/fakes + TDD**, **widget testing** (pump/finders/interactions) and **golden tests** (pixel snapshots), **integration/E2E testing** (`integration_test`, driving the real app, `patrol`), and how it all comes together — **testing each architectural layer + coverage + CI integration**. It's the safety net that makes every other module's architecture, refactoring, and scaling actually safe (Module 40–47, Module 50).

Why this module exists

Untested code can't be safely changed — every refactor, feature, or dependency bump risks silent regressions, so teams either freeze or ship bugs. Tests are what let you **move fast without breaking things**: they document behavior, catch regressions, enable confident refactoring, and gate CI. But testing done wrong (all slow E2E tests, brittle assertions, no isolation) is worse than useless. This module teaches the *right shape* (the pyramid), the tools per layer, and the discipline (isolation, fakes, TDD) that make tests fast, reliable, and valuable — the payoff the whole handbook's testable architecture was built for.

Module map

#	File	Topic	Level
1	01_testing_fundamentals_and_pyramid.md	Why test, the testing pyramid, test types, what to test	
2	02_unit_and_mocking.md	Unit tests, mocks/fakes (mocktail), TDD	
3	03_widget_and_golden_tests.md	Widget tests (pump/finders/interactions), golden tests	
4	04_integration_and_e2e.md	<code>integration_test</code> , E2E, driving the app, <code>patrol</code>	
5	05_testing_integration.md	Capstone: test each layer + coverage + CI	

Cross-references: Architecture (testable layers): 40/43/44. CI (run tests): Module 50. Error handling (test failure paths): Module 38. Networking/storage (fakes at boundaries): Module 16/Module 15. Performance (perf tests): Module 21. Scalable apps (tests as regression safety): Module 47.

Prerequisites

40 Clean Architecture (testable layers), 43 MVVM (state-sequence tests), 38 Error Handling, 14 Dependency Injection (inject fakes).

What you'll be able to do after this module

- Explain the testing pyramid and choose the right test type per concern.
- Write fast, isolated unit tests with mocks/fakes; practice TDD.
- Write widget tests (pump/finders/interactions) and golden tests.
- Write integration/E2E tests that drive the real app.
- Test each architectural layer, measure coverage meaningfully, and run tests in CI.

Capstone

Layered test suite: For a feature slice (domain/data/presentation), write unit tests (use case with fake repo; data repo with fake sources; view-model state sequences), widget tests (screen states + interactions), a golden test, and one integration test driving the flow — organized by the pyramid, with meaningful coverage and a CI job (`melos / flutter test + integration_test`) gating merges.

Summary

Testing = the right shape (pyramid: many unit, some widget, few E2E) + the right tools per layer (unit + mocks/fakes, widget + golden, integration/E2E) + discipline (isolation, fakes, TDD). It's what makes the handbook's testable architecture pay off — enabling confident refactoring, regression safety, and CI gating.

Testing Fundamentals & the Pyramid

Tests exist to let you **change code with confidence** — they document behavior, catch regressions, and gate CI. The **testing pyramid** prescribes the right *mix*: **many fast, isolated unit tests** (the base), **fewer widget tests** (the middle), and **few slow, broad integration/E2E tests** (the tip). Flutter maps this to **unit tests** (pure Dart logic — fastest), **widget tests** (a widget in a test harness, no device — medium), and **integration tests** (the real app on a device/emulator — slowest). The goal: **maximize confidence per unit of speed/maintenance** by pushing logic down to fast unit tests and reserving slow E2E for critical end-to-end flows.

Introduction

This file establishes *why* to test, the pyramid and its Flutter mapping, what each test type covers, and what to test (behavior, not implementation). It's the strategic frame for the tool-specific files.

Why this concept exists

Two failure modes plague testing: **too few tests** (can't refactor/ship safely) and **the wrong mix** (all slow, brittle E2E tests → slow CI, flaky, expensive to maintain). The pyramid encodes decades of experience: most confidence should come from **fast, isolated** tests, with expensive E2E used sparingly for critical paths. Getting the *shape* right is more important than raw count.

Real-world analogy

Testing a car: you **bench-test individual components** cheaply and exhaustively (unit — spark plugs, sensors), **test subsystems** on a rig (widget — the dashboard cluster), and only **road-test the whole car** occasionally (E2E — slow, expensive, but validates the real experience). A factory that *only* road-tests every car (all E2E) is slow and can't pinpoint faults; one that never assembles and drives one (no E2E) ships cars that don't actually work. The **pyramid** balances them.

Internal Working

Integration / E2E tests — FEW:
slow, broad, real device, critical
flows

maximize confidence per unit of
speed/maintenance -> push logic
down to unit

Widget tests — SOME: medium,
widget in harness, no device

Unit tests — MANY: fast, isolated,
pure Dart logic

- **Why test** (the payoff):
 - **Regression safety**: catch breakage when changing code (the #1 reason).
 - **Confident refactoring**: restructure freely with tests as the safety net (Module 47).
 - **Living documentation**: tests describe intended behavior.
 - **CI gating**: tests block bad merges (Module 50).
 - **Design pressure**: hard-to-test code reveals bad design (untestable = tightly coupled).
- **The pyramid (the right mix)**:
 - **Base — many unit tests**: fast (ms), isolated, deterministic; test **logic** (use cases, view models, repositories, pure functions). The bulk of your tests + confidence.
 - **Middle — some widget tests**: test **UI behavior** (rendering per state, interactions) in a test harness (no device); slower than unit, faster than E2E.
 - **Tip — few integration/E2E tests**: test **critical end-to-end flows** on a real device/emulator (login → browse → checkout); slow, broad, higher-maintenance — use **sparingly** for what unit/widget can't cover.
 - **Anti-shapes**: **ice-cream cone** (mostly E2E — slow/flaky/expensive) and **hourglass** (E2E + unit, no widget). Aim for the pyramid.
- **Flutter test types (map to pyramid)**:
 - **Unit test** (`package:test` / `flutter_test`): pure Dart, no `WidgetsFlutterBinding` (unless using `flutter_test` matchers); fastest.

- **Widget test** (`flutter_test` + `testWidgets`): pumps a widget tree in a test binding, no device; medium.
- **Integration test** (`integration_test`): runs the real app on device/emulator; slowest.
- **Golden test** (widget-level): pixel-snapshot comparison of rendered UI (`03_widget_and_golden_tests.md`).
- **What to test (behavior, not implementation)**: assert **observable behavior/outputs** (given input → expected output/state/render), **not** private internals or exact call sequences beyond what matters — so tests survive refactors. Cover **happy + unhappy paths** (errors, empty, edge cases — Module 38).
- **What NOT to over-test**: trivial getters, framework code, generated code; and don't chase 100% coverage for its own sake (`05_testing_integration.md`).
- **Testable architecture enables the pyramid**: Clean/MVVM (pure domain, view-model state, repository interfaces + fakes) lets most logic be tested by **fast unit tests** — the payoff of the architecture band (Module 40/Module 43).

Memory Representation

Not runtime — a **test suite shape**: a large base of unit tests, a medium band of widget tests, a small tip of E2E. The mental metric: **confidence per (speed × maintenance cost)** — favor fast, isolated tests.

Compiler Behavior

Tests are code; unit tests compile against pure Dart/interfaces (fakes injected), enabling isolation. Test files live in `test/` (unit/widget) and `integration_test/` (E2E).

Runtime Behavior

Unit tests run in the Dart VM (ms); widget tests in the Flutter test binding (no real device); integration tests on a device/emulator (seconds+). Speed differences are why the pyramid is shaped as it is.

Flutter Engine Behavior

Widget tests use a **test binding** (fake engine — no real rendering to a screen, controllable time via `pump`); integration tests use the real engine on-device. Unit tests don't touch the engine.

Dart VM Behavior

Unit tests are the fast tier (no binding/IO); parallelizable in CI. Integration tests are serial + slow (real app lifecycle).

Examples

```
// UNIT (fastest) – pure logic, no device; assert behavior (given -> expected)
test('total = sum of line subtotals', () {
  final order = Order('1', [LineItem('p', 2, 500), LineItem('q', 1, 300)]);
  expect(order.total, Money(1300, 'USD')); // behavior, not internals
});

// WIDGET (medium) – widget in a test harness, no device
testWidgets('shows error + retry on failure state', (tester) async {
  await tester.pumpWidget(wrap(ProfileView(state: ProfileError('Failed'))));
  expect(find.text('Failed'), findsOneWidget);
  expect(find.text('Retry'), findsOneWidget);
});

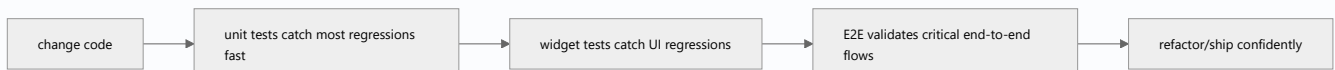
// INTEGRATION/E2E (slow, few) – real app, critical flow only
// integration_test: launch app -> login -> see home -> logout
```

Aim for the PYRAMID, avoid the ICE-CREAM CONE:

GOOD (pyramid): ~70% unit, ~20% widget, ~10% integration/E2E

BAD (cone): mostly E2E -> slow CI, flaky, hard to pinpoint failures

Diagrams



Common Mistakes

Mistake	Why it's wrong	Fix
Too few tests	Can't refactor/ship safely	Build the pyramid (unit-heavy)
Ice-cream cone (mostly E2E)	Slow, flaky, expensive, imprecise	Push logic to fast unit tests
Testing implementation details	Brittle; breaks on refactor	Test observable behavior/outputs
Only happy paths	Bugs live in errors/edges	Cover unhappy paths too
No isolation (real I/O in unit tests)	Slow, flaky, nondeterministic	Fakes/mocks at boundaries
Chasing 100% coverage	Wastes effort, false confidence	Meaningful coverage of behavior
Over-testing trivia/framework	Low value, maintenance	Skip trivial getters/generated code

Best Practices

- Build the **pyramid**: **many fast isolated unit tests** (logic), **some widget tests** (UI behavior), **few integration/E2E** (critical flows) — maximize confidence per speed/maintenance.
- **Test observable behavior** (given input → expected output/state/render), **not implementation details**; cover **happy + unhappy paths** (Module 38).
- Keep unit tests **isolated + deterministic** (fakes/mocks at boundaries — 02_unit_and_mocking.md); leverage **testable architecture** (Clean/MVVM) to push logic into fast unit tests.
- Avoid the **ice-cream cone**; don't over-test trivia/framework or chase 100% coverage; run tests in **CI** (Module 50).

Performance

Test *suite* speed is the point: unit tests (ms) give fast feedback + CI; E2E (seconds+) is the slow tail. A pyramid keeps CI fast + reliable; a cone makes CI slow + flaky. Push logic down to keep feedback loops tight (Module 47).

Advantages / Disadvantages

- + Regression safety, confident refactoring, living docs, CI gating, design pressure; fast feedback (pyramid).
- – Writing/maintaining tests costs effort; wrong shape (cone) is slow/flaky; over-testing wastes time; requires testable design.

Interview Questions

1. 🟢 **Why write tests?** — Regression safety, confident refactoring, living documentation, CI gating, and design pressure (untestable code signals bad design).
2. 🟢 **What is the testing pyramid?** — Many fast unit tests (base), fewer widget tests (middle), few slow integration/E2E tests (tip) — maximize confidence per speed/maintenance.
3. 🟡 **How does the pyramid map to Flutter test types?** — Unit (pure Dart logic, fastest), widget (`testWidgets` in a test binding, no device, medium), integration (`integration_test` on device, slowest).
4. 🟡 **Why is the ice-cream cone an anti-pattern?** — Mostly-E2E suites are slow, flaky, expensive, and can't pinpoint failures; push logic to fast unit tests.
5. 🟡 **Should you test implementation details?** — No — test observable behavior/outputs so tests survive refactors; avoid asserting private internals.
6. 🔴 **What should you NOT over-test, and why not chase 100% coverage?** — Trivial getters/framework/generated code; coverage % doesn't equal confidence — aim for meaningful behavior coverage including unhappy paths.
7. 🔴 **How does architecture enable the pyramid?** — Clean/MVVM (pure domain, view-model state, repository interfaces + fakes) lets most logic be covered by fast unit tests.

Senior Engineer Tips

- Push logic down: make it testable by a fast unit test (pure domain/view-model), and reserve E2E for the two or three flows that truly must work end-to-end — that shape keeps CI fast and reliable.
- Assert behavior, not internals; a test that breaks every refactor is testing the wrong thing and will get deleted or ignored.
- Treat "this is hard to test" as a design smell, not a testing problem — usually a missing seam (interface/DI) you should add.

Architect Perspective

The pyramid is the strategic shape that makes testing sustainable: fast, isolated unit tests carry most of the confidence, widget tests guard UI behavior, and a thin E2E tip validates critical journeys. It's the direct payoff of the architecture band — testable layers mean logic lands in the fast base — and the foundation for CI gating and confident evolution. Getting the shape and the behavior-not-implementation discipline right matters more than any single tool (02_unit_and_mocking.md, Module 40, Module 50).

Summary

- Tests enable confident change (regression safety, refactoring, docs, CI gating, design pressure).
- The pyramid: many fast unit tests, some widget tests, few integration/E2E — Flutter: unit (pure Dart) / widget (`testWidgets`) / integration (`integration_test`).
- Test behavior not implementation, cover unhappy paths, keep unit tests isolated/deterministic, avoid the ice-cream cone and 100%-coverage chasing.

Revision Notes

- Why: regression safety, confident refactor, docs, CI gate, design pressure. Pyramid: many unit (fast/isolated/logic) > some widget (UI behavior, `testWidgets`, no device) > few integration/E2E (`integration_test`, device, critical flows).
- Flutter types: unit (pure Dart, VM), widget (test binding, `pump` /finders), integration (real app/device), golden (pixel snapshot).
- Test behavior not internals; happy + unhappy paths; isolate unit (fakes/mocks); avoid ice-cream cone; don't over-test trivia / chase 100% coverage; architecture (Clean/MVVM) enables the fast base.

Practice Questions

1. What are the three pyramid layers and their Flutter test types?
2. Why is testing behavior better than testing implementation?
3. Why is an all-E2E (ice-cream cone) suite bad?

Coding Questions

1. Classify a set of concerns into unit/widget/integration tests.
2. Write a behavior-focused unit test (given input → expected output).
3. Identify an implementation-detail test and rewrite it to assert behavior.

Mini Project

Test strategy (Flutter): For a feature, produce a testing plan mapping concerns to pyramid layers (which are unit/widget/integration), justify the mix (unit-heavy), and write one behavior-focused unit test + one widget test stub + identify the single critical flow worth an E2E test. Acceptance: pyramid-shaped plan (unit-heavy, thin E2E); concerns mapped to the right type; behavior-not-implementation assertions; unhappy paths considered; one justified critical E2E flow; ice-cream cone avoided.

Unit Testing, Mocking & TDD

Unit tests verify a **single unit** (a function/class) in **isolation**, fast and deterministic — which requires **replacing its dependencies** with **test doubles: fakes** (working lightweight implementations — preferred), **mocks** (record/verify interactions — `mocktail` / `mockito`), and **stubs** (canned return values). The structure is **Arrange-Act-Assert** (`test('...', () { setup; act; expect })`), with `group` for organization and `setUp` / `tearDown` for shared fixtures. **TDD** (red → green → refactor) writes the test first to drive design. The golden rules: **isolate** (no real I/O), assert **behavior**, and **prefer fakes over mocks** where practical (less brittle).

Introduction

This file covers writing fast, isolated unit tests: structure (AAA, groups, setup), test doubles (fakes/mocks/stubs) and when to use each, mocktail basics, and the TDD cycle. It's the base of the pyramid (01_testing_fundamentals_and_pyramid.md) — where most tests and confidence live.

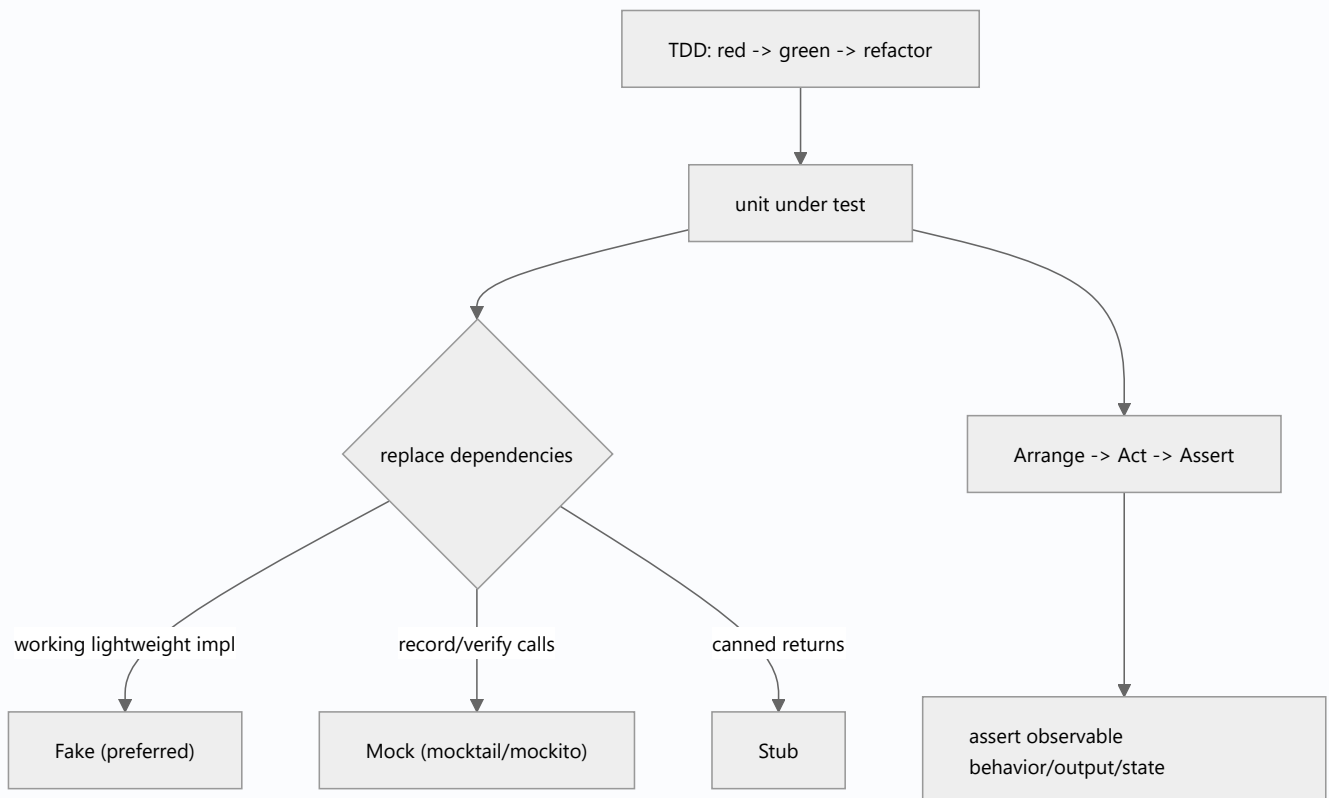
Why this concept exists

To test a unit in isolation you must remove its real dependencies (network, DB, other classes) — otherwise tests are slow, flaky, and non-deterministic, and a failure doesn't pinpoint the unit. Test doubles provide controlled stand-ins so you can drive the unit through every scenario (success/failure/edge) fast and deterministically. TDD uses tests to shape the code before writing it.

Real-world analogy

Unit testing with doubles is **bench-testing a car's ECU with a signal simulator**: you don't hook it to a real running engine (slow, dangerous, non-deterministic) — you feed it **simulated sensor inputs** (fakes/stubs) and verify its outputs, and sometimes check it **sent the right command** to a simulated actuator (mock). TDD is **writing the acceptance criteria before building the part**, so the part is built to pass.

Internal Working



- **Unit test structure (AAA):** **Arrange** (set up the unit + doubles + inputs), **Act** (call the method), **Assert** (`expect(actual, matcher)`). Use `group('...', () {...})` to organize and `setUp / tearDown` for shared fixtures. `flutter_test / package:test` provide `test`, `expect`, and rich **matchers** (`equals`, `isA<T>`, `throwsA`, `contains`, `predicate`).
- **Async tests:** mark the callback `async` and `await`; use `expectLater(future, completion(...)) / emitsInOrder([...])` for futures/streams.
- **Test doubles:**
 - **Fake (preferred):** a **working, lightweight implementation** (an in-memory repository, a fixed `Clock`). Behaves like the real thing for the test, isn't brittle, and often reusable. Best default.
 - **Stub:** returns **canned values** for calls (no verification) — for driving scenarios.
 - **Mock:** records calls and lets you **verify interactions** (`verify(() => x.foo(any())).called(1)`) and stub returns (`when(() => x.foo()).thenAnswer(...)`) — use when the **interaction itself is the behavior** (e.g., "the repository's save was called").
 - **Dummy:** passed but unused.
 - **Prefer fakes/stubs over mocks** where practical — mock-heavy tests couple to implementation (which methods were called) and become brittle; verify interactions only when they *are* the contract.

- **mocktail** (null-safe, no codegen — common choice): `class MockRepo extends Mock implements Repo {} ; when(() => repo.get(any())).thenAnswer((_) async => Success(x)); ; verify(() => repo.save(any())).called(1); ; registerFallbackValue(...)` for custom `any()` types. (`mockito` is similar but needs codegen.)
- **Isolation is mandatory: no real network/DB/file/time** in unit tests — inject fakes at the boundaries (repository interfaces, `clock`, etc.), which the architecture provides (Module 40/Module 14). Isolated → fast, deterministic, parallelizable.
- **TDD (red-green-refactor)**: (1) **red** — write a failing test for the desired behavior; (2) **green** — write the minimum code to pass; (3) **refactor** — clean up with the test as a safety net. Drives simple, testable design and full coverage of intended behavior. Use it where design is unclear/logic is important; pragmatic (not dogmatic) — not every line needs strict TDD.
- **Cover scenarios**: success, failure/errors, empty, boundary/edge cases — the unhappy paths where bugs live (Module 38).
- **Good tests: fast, isolated, deterministic, readable, one behavior per test**, descriptive names (`emits error when repository fails`).

Memory Representation

Test doubles are objects: fakes hold in-memory state, mocks record call events + stubbed responses, stubs hold canned returns. The test holds the unit + its doubles; assertions inspect returned values/state (or, for mocks, recorded calls). No real I/O objects.

Compiler Behavior

Mocks/fakes implement the dependency's interface (compile-checked). `mocktail` needs no codegen; `mockito` generates mock classes (`build_runner`). Injecting fakes requires the unit to depend on **abstractions** (interfaces) — the DIP payoff.

Runtime Behavior

Unit tests run in the Dart VM in milliseconds, deterministically (no I/O/time/randomness unless faked). Async tests await futures/streams. Failures pinpoint the unit.

Flutter Engine Behavior

Not applicable — pure unit tests don't touch the engine (that's widget/integration tests).

Dart VM Behavior

Fastest test tier (no binding); parallelizable in CI. Fakes/mocks are cheap objects.

Examples

```
import 'package:mocktail/mocktail.dart';
import 'package:flutter_test/flutter_test.dart';

// FAKE (preferred) – a working in-memory repository
class FakeOrderRepo implements OrderRepository {
  final _store = <String, Order>{};

  @override Future<Order?> findById(String id) async => _store[id];
  @override Future<void> save(Order o) async => _store[o.id] = o;
}

// MOCK (for interaction verification) – mocktail, no codegen
class MockOrderRepo extends Mock implements OrderRepository {}

void main() {
  group('PlaceOrder use case', () {
    late FakeOrderRepo repo;

    setUp(() => repo = FakeOrderRepo()); // shared fixture

    test('saves a placed order (fake: assert state)', () async { // AAA
      await repo.save(Order.draft('1')); // Arrange
      final uc = PlaceOrder(repo);
      final r = await uc('1'); // Act
      expect(r, isA<Success>()); // Assert (behavior)
      expect((await repo.findById('1'))!.status, OrderStatus.placed);
    });

    test('returns Failure on empty order (unhappy path)', () async {
      final uc = PlaceOrder(FakeOrderRepo());
      expect(await uc('missing'), isA<Failure>());
    });

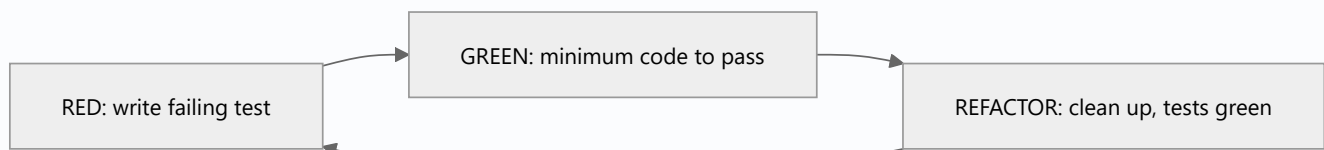
    test('mock: verifies save was called once (interaction is the contract)', () async {
      final mock = MockOrderRepo();
      when(() => mock.findById(any())).thenAnswer((_) async => Order.draftWithLines('1'));
      when(() => mock.save(any())).thenAnswer((_) async {});
      await PlaceOrder(mock)('1');
      verify(() => mock.save(any())).called(1); // verify interaction
    });
  });
}
```

```
});  
}
```

TDD cycle (red -> green -> refactor):

1. RED: write `test('applies 10% discount over $100', ...)` -> fails (no code)
2. GREEN: implement the minimum discount logic -> test passes
3. REFACTOR: clean up the code; tests stay green (safety net)

Diagrams



Common Mistakes

Mistake	Why it's wrong	Fix
Real network/DB/time in unit tests	Slow, flaky, non-deterministic	Inject fakes at boundaries
Mock-heavy tests (verify everything)	Brittle; couples to implementation	Prefer fakes; verify only true-contract interactions
Testing implementation details	Breaks on refactor	Assert behavior/output/state
Only happy-path tests	Bugs live in errors/edges	Cover failure/empty/boundary
Multiple behaviors per test	Hard to diagnose	One behavior per test
Shared mutable state across tests	Flaky, order-dependent	Fresh fixtures in <code>setUp</code>
Vague test names	Poor docs	Descriptive (<code>emits error when repo fails</code>)
Forgetting async await	False passes	<code>async / await</code> , <code>expectLater</code>

Best Practices

- Structure tests **AAA** (Arrange-Act-Assert), one **behavior per test**, descriptive names, `group` + `setUp` / `tearDown` for fixtures; use rich **matchers** and `expectLater` / `emitsInOrder` for async/streams.
- **Isolate**: no real I/O/time/randomness — inject **fakes** at boundaries (repository interfaces, `Clock`); **prefer fakes/stubs over mocks**, verifying interactions only when the interaction **is** the contract.
- Cover **happy + unhappy paths** (success/failure/empty/edge); assert **behavior, not internals**; keep tests fast + deterministic (parallelizable in CI).

- Use **TDD (red-green-refactor)** to drive design where logic/design is important — pragmatically, not dogmatically.

Performance

Unit tests are the fast tier (ms) and parallelize in CI — the whole point of isolation. Fakes/mocks are cheap. Real I/O in "unit" tests is the classic slowdown/flakiness source; keep it out. Fast unit tests = tight feedback loops (Module 47).

Advantages / Disadvantages

- + Fast, isolated, deterministic, precise-failure, parallelizable; exhaustive scenario coverage; TDD-driven design.
- – Requires testable design (interfaces/DI); mock overuse → brittleness; writing/maintaining tests costs effort; TDD has a learning curve.

Interview Questions

1. 🟢 **What makes a good unit test?** — Fast, isolated (no real I/O), deterministic, one behavior per test, descriptive name, asserting observable behavior.
2. 🟢 **Fake vs mock vs stub?** — Fake = working lightweight impl (preferred); stub = canned returns; mock = records/verifies interactions — use mocks only when the interaction is the contract.
3. 🟡 **Why prefer fakes over mocks?** — Mocks couple tests to implementation (which methods were called) → brittle; fakes behave like the real thing and survive refactors.
4. 🟡 **How do you isolate a unit that uses network/DB/time?** — Depend on abstractions (repository interface, `clock`) and inject fakes — the DIP/DI payoff.
5. 🟡 **What is the AAA structure?** — Arrange (setup + doubles + inputs), Act (call the method), Assert (`expect`).
6. 🔴 **Describe the TDD cycle and when to use it.** — Red (failing test) → green (minimum code) → refactor (clean up, tests green); use where logic/design matters, pragmatically.
7. 🔴 **What scenarios must unit tests cover beyond happy path?** — Failure/errors, empty, boundary/edge cases — where most bugs live.

Senior Engineer Tips

- Reach for a fake before a mock; write a small reusable in-memory fake for each repository interface, and only use mocks/ `verify` when "it called save once" is genuinely the behavior you're asserting.

- Keep unit tests hermetic — no real clock, network, DB, or randomness; inject a `clock` /fakes so tests are deterministic and parallel-safe.
- Name tests as behaviors and cover the unhappy paths first; those are where regressions actually happen, and good names make the suite double as documentation.

Architect Perspective

Unit tests + test doubles are the base of the pyramid and the direct dividend of a testable architecture: because domain/data/presentation depend on abstractions, you inject fakes and drive every scenario fast and deterministically. Preferring fakes over mocks keeps tests behavior-focused and refactor-resilient; TDD uses them to shape design. This fast, isolated base is what makes CI quick, refactoring safe, and the whole pyramid work (01_testing_fundamentals_and_pyramid.md, Module 40, Module 43).

Summary

- Unit tests verify one unit in isolation (fast/deterministic) via AAA, using test doubles — prefer fakes (working impls) over mocks (interaction verification), stubs for canned returns.
- Isolate (no real I/O/time) by injecting fakes at abstraction boundaries; assert behavior; cover happy + unhappy paths; use mocktail (no codegen).
- TDD (red-green-refactor) drives design pragmatically; keep tests fast, deterministic, one-behavior, well-named.

Revision Notes

- Structure: AAA + `group` / `setUp` / `tearDown` ; matchers (`isA` , `throwsA` , `equals`), `async` `await` / `expectLater` / `emitsInOrder` .
- Doubles: fake (working impl, preferred) > stub (canned) > mock (verify interactions — mocktail: `when(...).thenAnswer` , `verify(...).called(n)` , `registerFallbackValue`). Prefer fakes; verify only true-contract interactions.
- Isolate (no real I/O/time → inject fakes at interfaces/ `clock`); assert behavior not internals; cover happy + unhappy; TDD red-green-refactor (pragmatic); fast/deterministic/parallelizable.

Practice Questions

1. When do you use a mock vs a fake?
2. How do you keep a unit test deterministic when the unit uses the network?
3. Walk through a TDD cycle for a small feature.

Coding Questions

1. Write AAA unit tests (success + failure) for a use case using a fake repo.
2. Use mocktail to verify a repository interaction where the interaction is the contract.
3. Implement a small feature via TDD (red-green-refactor).

Mini Project

Unit test suite + TDD (Flutter): For a use case + repository, write an in-memory **fake** repository and unit tests (AAA) covering success/failure/empty/edge, plus one **mock**-based test where verifying an interaction (`save` called once) is the contract, all isolated (no real I/O). Then TDD one new rule (red → green → refactor). Acceptance: fake-based isolated tests (no real I/O/time); AAA + descriptive names + one behavior each; happy + unhappy paths; mock used only for a true-contract interaction; a TDD'd rule with passing tests; fast + deterministic.

Widget & Golden Tests

Widget tests verify **UI behavior** without a device: `testWidgets` **pumps** a widget into a fake test binding, you locate elements with **finders** (`find.text` / `byType` / `byKey`), drive **interactions** (`tester.tap` / `enterText` + `pump` / `pumpAndSettle`), and **assert** with matchers (`findsOneWidget`).

Golden tests go further — they render a widget and compare it **pixel-for-pixel** against a saved reference image (`matchesGoldenFile`), catching *visual* regressions (layout/color/spacing) that behavioral assertions miss. Both are the pyramid's middle: slower than unit, far faster than E2E, and the way to test screens/widgets and their states/interactions.

Introduction

This file covers widget testing (`pump`/`finders`/`interactions`/`pumpWidget` harness, faking dependencies, `pump` vs `pumpAndSettle`) and golden (snapshot) testing (when/how, managing goldens). It's the UI-behavior middle of the pyramid (`01_testing_fundamentals_and_pyramid.md`).

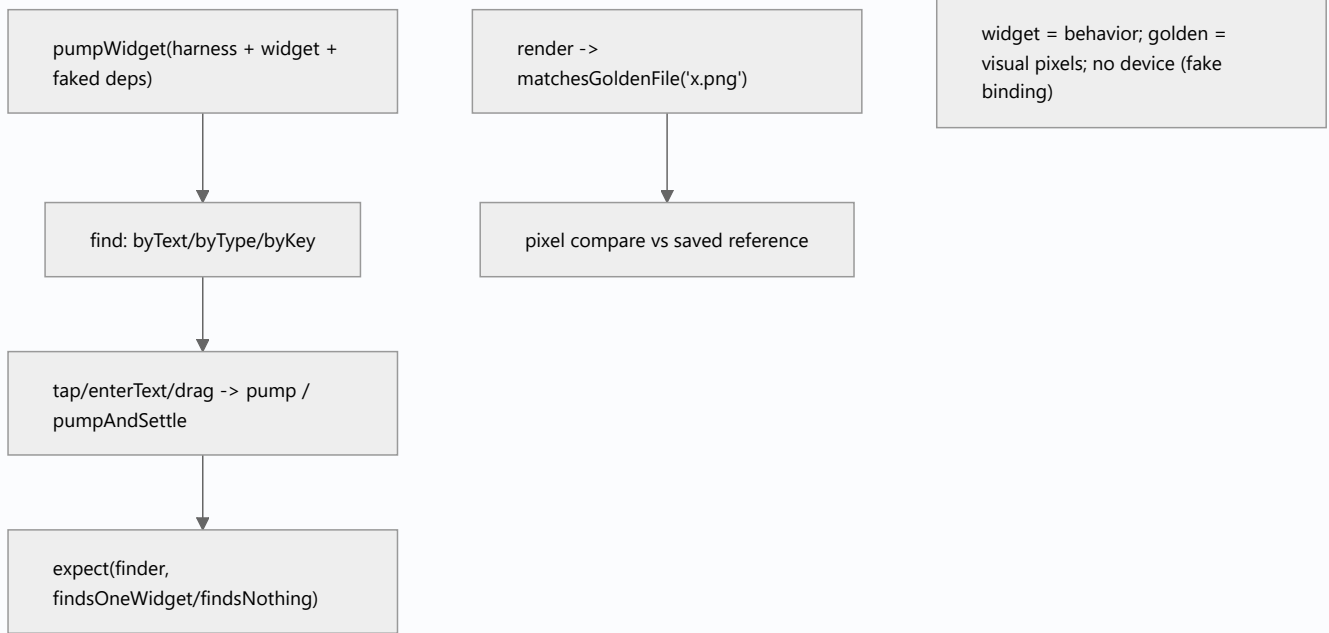
Why this concept exists

Unit tests can't verify that the UI **renders correctly and responds to interaction**; E2E tests can but are slow/flaky. Widget tests fill the gap: real widget rendering in a controllable, device-free binding — fast enough to run many, precise enough to catch UI regressions. Golden tests add **visual** verification (pixels), catching layout/style regressions that "find the text" assertions can't.

Real-world analogy

A widget test is **testing a dashboard cluster on a bench rig**: you power it with simulated inputs (faked state), **press its buttons** (tap), and **check the gauges read correctly** (finders/matchers) — without installing it in a car (device). A golden test is **photographing the lit dashboard and comparing it to an approved photo** — catching a misaligned needle or wrong-colored light that "the speed value is 60" wouldn't.

Internal Working



- **testWidgets + WidgetTester**: the entry point. `await tester.pumpWidget(widget)` builds a tree in the **test binding** (a fake engine — no real screen, controllable time). Wrap the widget under test in the needed ancestors (`MaterialApp` / `Directionality` / providers) — a **test harness**.
- **Finders** (locate elements): `find.text('Retry')`, `find.byType(ElevatedButton)`, `find.byKey(Key('submit'))`, `find.byIcon`, `find.byWidgetPredicate`. Prefer **byKey / byType** for stability over brittle text where appropriate.
- **Interactions**: `await tester.tap(finder)`, `enterText(finder, 'x')`, `drag`, `fling`, `longPress` — each followed by a **pump** to process the frame.
- **pump vs pumpAndSettle**: `pump()` advances **one frame** (or a duration); `pumpAndSettle()` pumps **until no frames are scheduled** (animations/async settle). Use `pumpAndSettle` after navigation/animations; but it **hangs on infinite animations** — use `pump(duration)` there.
- **Matchers**: `findsOneWidget`, `findsNothing`, `findsNWidgets(n)`, `findsWidgets`; combine with `expect`.
- **Faking dependencies**: inject **fakes** (fake view model/use case/repository) via the harness (Provider override, constructor, DI) so the widget test is **isolated + deterministic** — no real network (02_unit_and_mocking.md). Test **each UI state** (loading/data/empty/error) by feeding the corresponding state/fake (Module 43).
- **What widget tests verify**: rendering per state, presence/absence of elements, interaction outcomes (tap → state change → new render), form validation, navigation triggers — **UI behavior**, not business logic (that's unit tests).
- **Golden (snapshot) tests**: `await expectLater(find.byType(MyWidget), matchesGoldenFile('goldens/my_widget.png'))`. First run (`--update-goldens`) **saves** the reference; later

runs **compare pixel-for-pixel** and fail on diffs. Catches **visual** regressions (layout/spacing/color/font).

- **Manage carefully:** goldens are **platform/font-sensitive** (render can differ across OS/CI) — generate on a **consistent environment** (CI or `golden_toolkit` / `alchemist` for device-independent rendering + multi-device goldens), **review golden diffs** as artifacts, and **regenerate deliberately** on intended UI changes. Don't golden-test volatile/dynamic content (timestamps, random data).
- Use goldens for **stable, design-critical** widgets (design-system components, key screens), not everything.
- **Test keys:** add `key` s to widgets you need to find reliably (avoids brittle text-based finders), especially for lists/dynamic UI.
- **Speed/reliability:** widget tests run in the VM's test binding (no device) — fast enough for many; keep them **isolated** (fakes) and **deterministic** (control time via `pump` , avoid real async).

Memory Representation

Widget tests build a real widget/element/render tree in the test binding (in-memory, no GPU). Goldens are saved PNG references compared against freshly-rendered pixels. Fakes provide deterministic state.

Compiler Behavior

Widget tests compile against real widgets + fakes; goldens reference saved image files. Test keys/finders are compile-time widget references.

Runtime Behavior

The test binding renders synchronously under your control (`pump`); interactions + async resolve only when you pump. Golden comparison is pixel-diff (fails on any mismatch beyond tolerance). No real device/vsync.

Flutter Engine Behavior

The **test binding fakes the engine**: layout/paint happen into an offscreen surface, time is manually advanced (`pump`), and there's no real rasterization to a screen — enabling fast, deterministic UI tests. Goldens capture that offscreen render.

Dart VM Behavior

Runs in the VM with the Flutter test binding — slower than pure unit (tree building) but far faster than device E2E; parallelizable.

Examples

```
import 'package:flutter_test/flutter_test.dart';

// Harness: wrap the widget with needed ancestors + faked state
Widget harness(Widget child) => MaterialApp(home: Scaffold(body: child));

void main() {
  testWidgets('error state shows message + retry, tap triggers retry', (tester) async {
    var retried = false;
    await tester.pumpWidget(harness( // pump into fake binding
      ProfileView(state: const ProfileError('Failed'), onRetry: () => retried = true),
    ));
    expect(find.text('Failed'), findsOneWidget); // finder + matcher
    expect(find.text('Retry'), findsOneWidget);
    await tester.tap(find.text('Retry')); // interaction
    await tester.pump(); // process the frame
    expect(retried, isTrue); // behavior assertion
  });

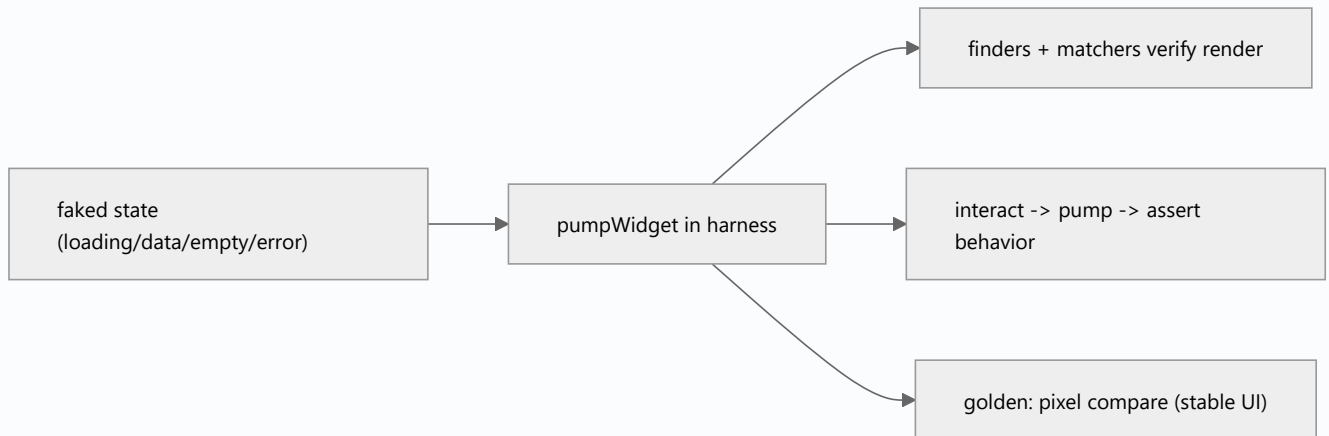
  testWidgets('loading state shows a spinner', (tester) async {
    await tester.pumpWidget(harness(const ProfileView(state: ProfileLoading())));
    expect(find.byType(CircularProgressIndicator), findsOneWidget);
  });

  // GOLDEN – visual snapshot of a design-system component (stable UI)
  testWidgets('AppButton matches golden', (tester) async {
    await tester.pumpWidget(harness(const AppButton(label: 'Save')));
    await expectLater(find.byType(AppButton), matchesGoldenFile('goldens/app_button.png'));
    // First run: flutter test --update-goldens (saves reference). Later: pixel compare.
  });
}
```

pump vs pumpAndSettle:

```
await tester.pump();           // advance ONE frame (or pump(Duration))
await tester.pumpAndSettle();  // pump until no frames scheduled (animations/async settle)
                               // WARNING: hangs on infinite animations -> use pump(duration)
```

Diagrams



Common Mistakes

Mistake	Why it's wrong	Fix
Real network/deps in widget tests	Slow/flaky/non-deterministic	Inject fakes (fake VM/use case)
Forgetting to <code>pump</code> after interaction	UI not updated → wrong assertion	<code>pump</code> / <code>pumpAndSettle</code> after actions
<code>pumpAndSettle</code> with infinite animation	Test hangs	<code>pump(duration)</code> instead
Brittle text-only finders	Break on copy changes	Use <code>byKey</code> / <code>byType</code> where sensible
Golden-testing volatile content	Constant false failures	Golden only stable/design-critical UI
Goldens on inconsistent envs	Platform/font pixel diffs	Generate on consistent env / <code>golden_toolkit</code> / <code>alchemist</code>
Not testing all UI states	Missed empty/error regressions	Test loading/data/empty/error
Missing harness ancestors	Widget throws (no Directionality/Material)	Wrap in required ancestors

Best Practices

- Use `testWidgets` + a **harness** (needed ancestors) with **faked state/deps** to test **each UI state** (loading/data/empty/error) + interactions; assert with **finders + matchers**; `pump` / `pumpAndSettle` appropriately.
- Prefer `byKey` / `byType` finders for stability; add `Key` s to dynamic widgets; keep widget tests **isolated + deterministic** (no real I/O/time).
- Use **golden tests** for **stable, design-critical** widgets (design-system components, key screens); generate on a **consistent environment** (or `golden_toolkit` / `alchemist`), **review diffs**, and **regenerate deliberately**; avoid goldening volatile content.
- Test **UI behavior** here (rendering/interactions/navigation triggers), leaving **business logic to unit tests**.

Performance

Widget tests run in the test binding (no device) — fast enough for many, but slower than unit tests (tree building) → keep the pyramid unit-heavy. Golden pixel-diffs add a bit; run + parallelize in CI. `pumpAndSettle` on infinite animations is the classic hang (use `pump(duration)`).

Advantages / Disadvantages

- + Verifies real UI rendering + interactions device-free; per-state coverage; goldens catch visual regressions; fast enough for many.
- – Slower than unit; goldens are env/font-sensitive + need review/regeneration; brittle finders; harness/faking setup; `pumpAndSettle` pitfalls.

Interview Questions

1. 🟢 **What does a widget test verify, and how?** — UI behavior (rendering/interactions) device-free: `pumpWidget` into a test binding, locate with finders, drive interactions + `pump` , assert with matchers.
2. 🟢 **`pump` vs `pumpAndSettle` ?** — `pump` advances one frame (or a duration); `pumpAndSettle` pumps until no frames are scheduled — but hangs on infinite animations (use `pump(duration)`).
3. 🟡 **What is a golden test and what does it catch?** — A pixel snapshot compared against a saved reference (`matchesGoldenFile`); catches visual regressions (layout/spacing/color/font) that behavioral assertions miss.
4. 🟡 **How do you keep widget tests isolated/deterministic?** — Inject fakes for the view model/use case/repository via the harness; no real network/time; test each state by feeding it.
5. 🟡 **Why prefer `byKey` / `byType` over text finders?** — Text finders break on copy changes; keys/types are stable references (add `Key` s to dynamic widgets).

6. **What makes golden tests tricky, and how do you manage them?** — Platform/font pixel differences; generate on a consistent env (or `golden_toolkit` / `alchemist`), review diffs, regenerate deliberately, and avoid goldening volatile content.
7. **What belongs in a widget test vs a unit test?** — Widget: rendering/interactions/navigation triggers (UI behavior); unit: business logic — don't test logic through the widget.

Senior Engineer Tips

- Drive widget tests off faked state so you can test loading/data/empty/error deterministically; testing UI through real network is slow, flaky, and misses states.
- Reserve golden tests for stable design-system components and key screens, generate them on a controlled environment (CI or `alchemist` / `golden_toolkit`), and treat golden diffs as reviewable artifacts — casual goldens everywhere are a maintenance sink.
- Add `key` s and prefer type/key finders; text-based finders that break on every copy tweak are why teams start ignoring widget tests.

Architect Perspective

Widget and golden tests are the pyramid's UI-behavior middle: they verify that state renders correctly and interactions work (widget) and that the pixels stay right (golden), device-free and fast enough to run broadly. Driven off faked state (the MVVM payoff), they cover the states/interactions unit tests can't and the visuals E2E is too slow to guard — completing behavioral + visual regression safety before the thin E2E tip (Module 43, 04_integration_and_e2e.md, 01_testing_fundamentals_and_pyramid.md).

Summary

- Widget tests (`testWidgets`): pump a widget into a fake binding, find elements (finders), interact (tap/enterText + pump/pumpAndSettle), assert (matchers) — verifying UI behavior device-free, per state, with faked deps.
- Golden tests: pixel-snapshot comparison (`matchesGoldenFile`) catching visual regressions — for stable/design-critical UI, on a consistent env, reviewed + regenerated deliberately.
- Prefer key/type finders; isolate with fakes; `pumpAndSettle` (not on infinite animations); leave business logic to unit tests.

Revision Notes

- `testWidgets((tester){...}) : pumpWidget(harness) ; finders ( find.text/byType/byKey/byIcon/byWidgetPredicate ); interactions ( tap/enterText/drag +`

```
pump / pumpAndSettle ); matchers ( findsOneWidget/findsNothing/findsNWidgets ).
```

- `pump` = one frame/duration; `pumpAndSettle` = until idle (hangs on infinite animations → `pump(duration)`); wrap in required ancestors (harness); inject fakes → isolated/deterministic; test each state (loading/data/empty/error).
- Golden: `matchesGoldenFile` pixel compare; `--update-goldens` saves; env/font-sensitive → consistent env / `golden_toolkit` / `alchemist`; review diffs, regenerate deliberately, only stable/design-critical UI; prefer key/type finders; UI behavior here, logic in unit tests.

Practice Questions

1. How do you test each UI state of a screen device-free?
2. When do you use golden tests, and what makes them tricky?
3. `pump` vs `pumpAndSettle` — and when does the latter hang?

Coding Questions

1. Write widget tests for loading/data/empty/error states with a faked VM.
2. Test a tap → state-change → re-render interaction (with `pump`).
3. Add a golden test for a design-system component.

Mini Project

Widget + golden suite (Flutter): For a screen with loading/data/empty/error states, write widget tests (harness + faked view model) verifying each state's render + a tap interaction (retry → triggers reload), using key/type finders, plus a golden test for one stable design-system component (generated on a consistent env). Acceptance: each UI state tested with faked deps (isolated/deterministic); interaction verified with proper `pump`; stable finders (key/type); one golden for a design-critical widget (reviewable, deliberately regenerated); no business logic tested through the widget.

Integration & E2E Testing

Integration/E2E tests are the pyramid's **thin tip**: they run the **real app on a real device/emulator** (via the `integration_test` package), driving actual user flows end-to-end (launch → login → navigate → act → verify) across the whole stack — real widgets, real navigation, often a **test/staging backend** or seeded fakes. They give the **highest-fidelity confidence** ("the app actually works") but are **slow, broad, and flakier**, so you write **few** of them, only for **critical journeys**. `patrol` extends them to native interactions (permissions, notifications, webviews) `integration_test` can't reach.

Introduction

This file covers E2E/integration testing: `integration_test` setup, driving real flows, test data/backends, `patrol` for native interactions, and the discipline of keeping E2E few, focused, and reliable. It's the pyramid tip (01_testing_fundamentals_and_pyramid.md).

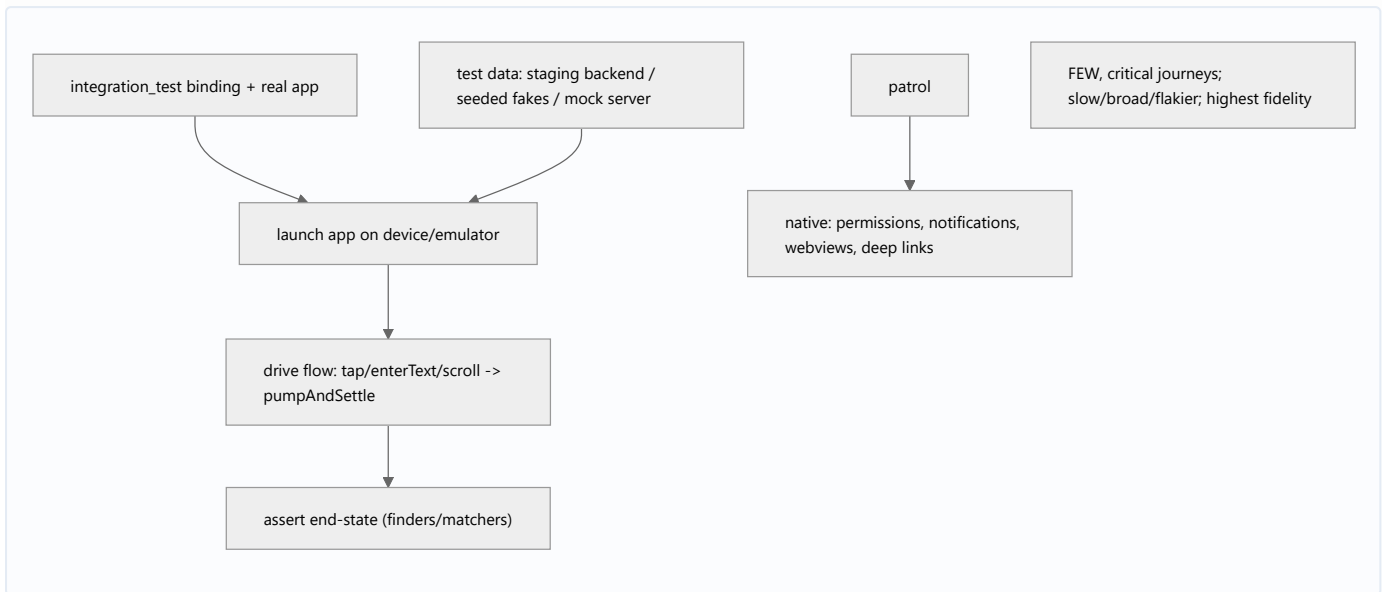
Why this concept exists

Unit + widget tests verify pieces in isolation but not that **the whole app works together** on a real device (real navigation, plugins, platform behavior). E2E tests validate the **actual end-to-end experience** — the ultimate confidence — but their cost (slow, flaky, maintenance) means they must be **rationed** to the flows that truly matter, sitting atop a solid unit/widget base.

Real-world analogy

E2E is the **full road test of the assembled car**: you actually **drive it through the real route** (start → highway → parking) to confirm the whole thing works as a driver experiences it. It's the most convincing test, but you don't road-test every bolt combination — you road-test the **key journeys**, relying on bench/rig tests (unit/widget) for the rest. Too many road tests = slow, expensive, and every pothole (flake) stalls the line.

Internal Working



- **integration_test package**: the official E2E tool. Tests live in `integration_test/`, use `IntegrationTestWidgetsFlutterBinding.ensureInitialized()`, and drive the **real app** (`app.main()` then interact) on a **device/emulator** — real widgets, navigation, plugins. Run with `flutter test integration_test/...` or on devices/ `flutter drive`; runs in CI on emulators/device farms (Firebase Test Lab).
- **Driving flows**: same `WidgetTester` API as widget tests (`tap` / `enterText` / `scrollUntilVisible` + `pumpAndSettle`), but against the **full running app** — so you traverse **multiple screens/features** (login → home → detail → action → verify), exercising the real stack.
- **Test data / backend** (a key decision):
 - **Staging/test backend**: highest fidelity, but slower/flakier + needs seeded, isolated data (reset between runs) to be deterministic.
 - **Seeded fakes / mock server**: swap in a fake network layer or a **local mock server** for determinism/speed while still exercising app integration — a common, more reliable choice.
 - **Reset state** between tests (fresh login/data) to avoid order-dependence + flakiness.
- **patrol** (extends `integration_test`): drives **native OS interactions** that `integration_test` can't — **granting/denying permissions**, tapping **system notifications**, interacting with **native webviews**, handling **native dialogs/deep links**, and more robust element selection. Use it when a critical flow crosses into native territory (e.g., "grant camera permission → capture").
- **Keep E2E few + focused (the discipline)**: only **critical user journeys** (onboarding/login, core purchase/checkout, primary feature happy path). E2E is **slow (seconds-minutes)**, **broad (fails don't pinpoint)**, and **flakier** (real async/network/timing) — so a small, curated set, not comprehensive coverage. Everything else → unit/widget.
- **Reliability practices** (E2E is flaky by nature): deterministic test data, `pumpAndSettle` / explicit waits for async (avoid arbitrary `sleep`), retry policy for known-flaky infra, stable finders (keys),

isolated/reset state, and running on consistent emulators. Quarantine + fix flaky tests fast (flaky E2E erodes trust).

- **What E2E validates that lower tests can't:** real navigation/routing, plugin/platform behavior, DI wiring end-to-end, cross-feature flows, and "does the assembled app actually work."
- **Performance/other E2E:** `integration_test` can also capture **performance traces** (frame timings) for perf regression checks (Module 21).

Memory Representation

E2E runs the full app process on a device/emulator (real memory/lifecycle). Test data lives in a staging backend or seeded fake/mock server. The test drives the real widget tree and asserts on real rendered elements.

Compiler Behavior

Integration tests compile the real app + test driver; `patrol` adds native test infrastructure. Tests reference real app entry points (`app.main()`).

Runtime Behavior

The real app launches + runs on-device; the test drives it via the binding; real async/network/navigation occur (hence slower + flakier). `pumpAndSettle` waits for real frames/async. Perf traces can be collected.

Flutter Engine Behavior

Uses the **real engine** on a real device/emulator (real rendering/vsync/plugins) — unlike widget tests' fake binding. This is what gives high fidelity and also the slowness/flakiness.

Dart VM Behavior

Runs the full app in AOT/JIT on-device; serial + slow relative to unit/widget; the reason to keep E2E few.

Examples

```
// integration_test/app_test.dart – drive the REAL app end-to-end
import 'package:integration_test/integration_test.dart';
import 'package:flutter_test/flutter_test.dart';
import 'package:myapp/main.dart' as app;

void main() {
  IntegrationTestWidgetsFlutterBinding.ensureInitialized();

  testWidgets('critical journey: login -> see home', (tester) async {
    app.main(); // launch the real app
    await tester.pumpAndSettle();

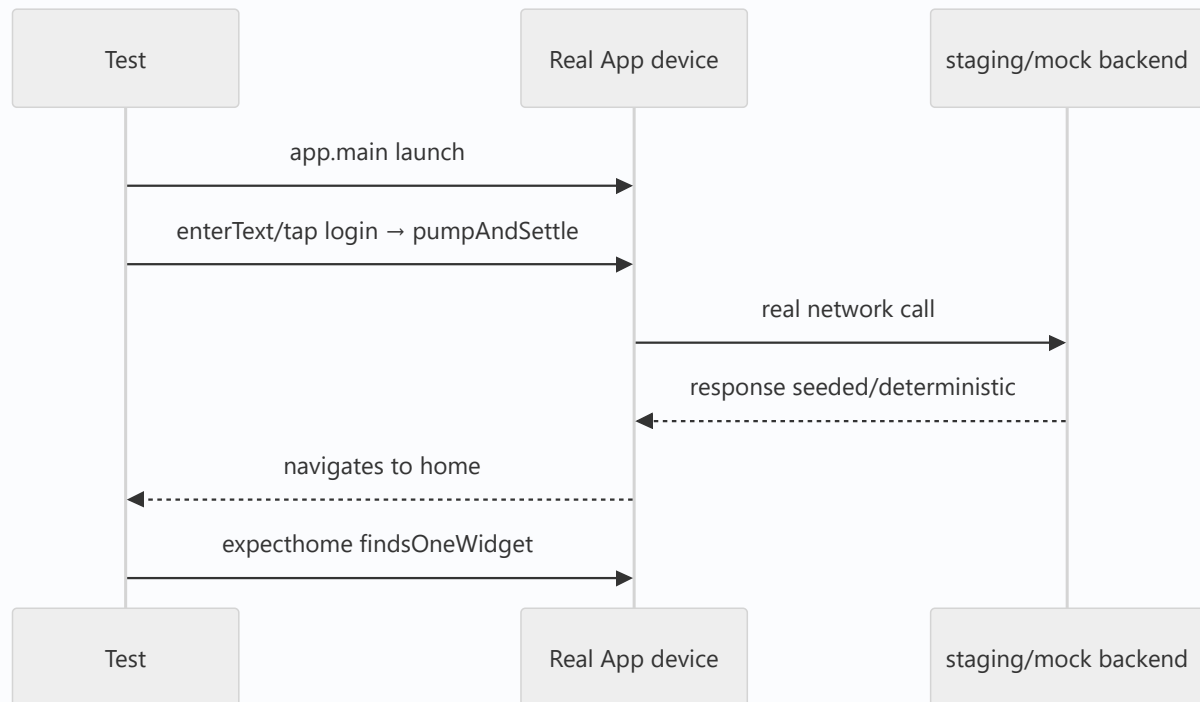
    await tester.enterText(find.byKey(const Key('email')), 'test@example.com');
    await tester.enterText(find.byKey(const Key('password')), 'secret');
    await tester.tap(find.byKey(const Key('loginBtn')));
    await tester.pumpAndSettle(); // wait for real navigation/async

    expect(find.byKey(const Key('homeScreen')), findsOneWidget); // end-state assertion
  });
}

// Backend: point at a seeded staging env OR a local mock server for determinism; reset state per run.
```

```
// patrol – native interactions integration_test can't do (permissions/notifications)
// patrolTest('grant camera permission then capture', ($) async {
//   await $.native.grantPermissionWhenInUse(); // native permission dialog
//   await $(#capture).tap();
//   expect($(#photoPreview), findsOneWidget);
// });
```

Diagrams



Common Mistakes

Mistake	Why it's wrong	Fix
Too many E2E tests	Slow CI, flaky, high maintenance	Few, critical journeys only; push down to unit/widget
Non-deterministic test data	Flaky/order-dependent	Seeded/reset staging or mock server
Arbitrary <code>sleep()</code> for async	Flaky + slow	<code>pumpAndSettle</code> /explicit waits
Text finders in E2E	Break on copy/localization	Stable <code>key</code> finders
Ignoring flaky E2E	Erodes trust in the suite	Quarantine + fix fast; retry known-infra flakes
Using E2E for logic coverage	Slow/imprecise	Unit tests for logic
Needing native interaction with plain <code>integration_test</code>	Can't grant permissions/tap notifications	Use <code>patrol</code>
No state reset between tests	Cross-test contamination	Reset app/backend state per test

Best Practices

- Keep E2E **few and focused** on **critical journeys** (onboarding/login, checkout, core happy path); rely on the **unit/widget base** for everything else.

- Use `integration_test` to drive the **real app** on device/emulator; make it **deterministic** (seeded/reset staging **or** a local mock server, stable `key` finders, `pumpAndSettle` /explicit waits — no `sleep`).
- Use `patrol` when a critical flow needs **native interactions** (permissions, notifications, webviews, deep links).
- **Manage flakiness aggressively** (reset state, consistent emulators, quarantine + fix); run in **CI** (device farm/emulator); optionally capture **perf traces**.

Performance

E2E is the slow tail (seconds-minutes/test, serial, real device) — the reason for its thin tip. In CI it runs on emulators/device farms, often in parallel across devices but still the pipeline's slow part. Keep the set minimal; push coverage down to fast unit/widget tests to keep CI quick (01_testing_fundamentals_and_pyramid.md).

Advantages / Disadvantages

- + Highest fidelity ("the app really works"), real navigation/plugins/platform, cross-feature flows, DI wiring validated end-to-end, native flows via patrol, perf traces.
- – Slow, broad (imprecise failures), flakier (real async/network/timing), high maintenance, needs test-data strategy + consistent devices; must be rationed.

Interview Questions

1. ● **What do integration/E2E tests verify, and where in the pyramid?** — That the assembled app works end-to-end on a real device (real navigation/plugins/stack); the thin tip — few tests for critical journeys.
2. ● **What tool does Flutter use for E2E?** — The `integration_test` package (drives the real app on device/emulator with the `WidgetTester` API); `patrol` extends it for native interactions.
3. ● **Why keep E2E tests few?** — They're slow, broad (imprecise failures), and flaky (real async/network/timing) with high maintenance — reserve for critical flows, push the rest to unit/widget.
4. ● **How do you make E2E deterministic?** — Seeded/reset staging data or a local mock server, stable `key` finders, `pumpAndSettle` /explicit waits (no `sleep`), consistent emulators, per-test state reset.
5. ● **When do you need `patrol` over `integration_test` ?** — For native OS interactions `integration_test` can't do: granting permissions, tapping system notifications, native webviews/dialogs/deep links.
6. ● **What does E2E validate that unit/widget tests can't?** — Real navigation/routing, plugin/platform behavior, end-to-end DI wiring, and cross-feature journeys on the real

engine.

7. 🟡 **How do you handle E2E flakiness?** — Deterministic data + waits + stable finders + consistent devices; quarantine and fix flaky tests fast (and retry only genuine infra flakes) — flaky E2E erodes trust.

Senior Engineer Tips

- Ration E2E to the handful of journeys that must never break (login, checkout, core flow); comprehensive E2E coverage is a slow, flaky trap — that's what the unit/widget base is for.
- Kill flakiness at the source: deterministic seeded/mock data, `pumpAndSettle` /explicit conditions instead of `sleep`, stable `key` finders, and consistent emulators; one ignored flaky test rots the whole suite's credibility.
- Reach for `patrol` only when a critical flow genuinely crosses into native (permissions/notifications); otherwise keep E2E in plain `integration_test`.

Architect Perspective

Integration/E2E is the confidence capstone of the pyramid: it validates that the assembled app — navigation, plugins, DI, cross-feature flows — actually works on a real device, which no isolated test can. Its cost mandates discipline: a thin, curated set of critical journeys, made deterministic and reliable, atop a broad fast base of unit/widget tests, with `patrol` for native reach. Run in CI, it's the final gate that "it really works" before release — the top of the testing strategy the whole handbook's testable architecture supports (05_testing_integration.md, Module 50, Module 51).

Summary

- E2E/integration = pyramid tip: `integration_test` drives the real app on a device/emulator through critical journeys (highest fidelity, slow/broad/flakier) — keep them few.
- Make deterministic (seeded/reset staging or mock server, stable `key` finders, `pumpAndSettle` /explicit waits, consistent emulators, per-test reset); manage flakiness aggressively.
- Use `patrol` for native interactions (permissions/notifications/webviews); run in CI; validates end-to-end navigation/plugins/DI that unit/widget can't.

Revision Notes

- `integration_test`: `IntegrationTestWidgetsFlutterBinding.ensureInitialized()`, `app.main()`, drive with `WidgetTester` across real screens on device/emulator; run via `flutter test integration_test` /device farm (Firebase Test Lab); can capture perf traces.

- Few + critical journeys only (slow/broad/flaky); deterministic (seeded/reset staging OR mock server, `Key` finders, `pumpAndSettle` /explicit waits — no `sleep` , consistent emulators, reset per test); quarantine/fix flakes.
- `patrol` for native (permissions/notifications/webviews/deep links); validates navigation/plugins/DI/cross-feature end-to-end; run in CI atop a fast unit/widget base.

Practice Questions

1. Why are E2E tests kept few, and what do they uniquely validate?
2. How do you make an E2E test deterministic?
3. When do you reach for `patrol` instead of `integration_test` ?

Coding Questions

1. Write an `integration_test` driving a login → home critical journey.
2. Make it deterministic (seeded/mock backend + stable finders + waits).
3. Sketch a `patrol` test granting a permission then using the feature.

Mini Project

E2E critical journey (Flutter): Write an `integration_test` driving one critical journey (e.g., login → browse → open detail → back), pointing at a seeded staging env or local mock server for determinism, using `Key` finders + `pumpAndSettle` (no `sleep`) with per-test state reset — and sketch a `patrol` test for a native-permission flow. Keep it to the single critical journey (rely on unit/widget for the rest). Acceptance: real-app E2E of one critical journey; deterministic (seeded/mock data, stable finders, proper waits, reset); flakiness managed; `patrol` used only for native; sits atop a fast unit/widget base (thin tip); runnable in CI.

Testing Integration (Capstone: Test Every Layer + Coverage + CI)

Bring the pyramid to a real feature: **unit-test each architectural layer** (domain use cases with fake repos; data repositories with fake sources; presentation view-models via state sequences), **widget-test** the screen's states + interactions (+ a golden for stable UI), and write **one integration test** for the critical journey — then **measure coverage meaningfully** (behavior + unhappy paths, not a 100% vanity metric) and **gate it all in CI** (`flutter test + integration_test` , or `melos run test` for a monorepo). The result: a **pyramid-shaped, layer-complete, CI-gated** suite that makes the feature safe to change — the payoff of the whole handbook's testable architecture.

Introduction

This module capstone composes unit/widget/golden/integration testing into one layered suite for a feature, adds coverage measurement and CI gating, and shows how testable architecture makes it straightforward. It's the practical "how it all fits" deliverable.

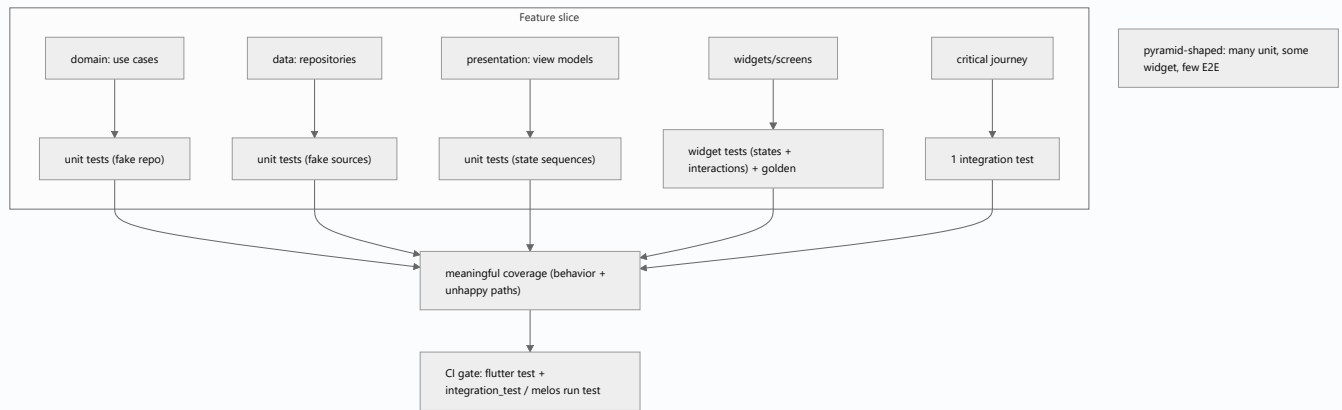
Why this concept exists

Individual test types are only valuable **assembled into a coherent, pyramid-shaped suite** that covers every layer, runs in CI, and is measured meaningfully. This capstone demonstrates that assembly — the difference between "we have some tests" and "the feature is safe to change."

Real-world analogy

It's the **full QA program for one car model**: bench tests for every component (unit per layer), rig tests for subsystems (widget/golden), a road test of the key journey (E2E), a **coverage report** of what's verified, and an **assembly-line gate** that won't ship a car that fails any of them (CI). Together they let the factory iterate the model confidently.

Internal Working



- **Test every layer (unit — the base):**

- **Domain:** use cases with **fake repositories** (assert `Result` /behavior + invariants) — pure Dart, fastest (Module 40/Module 46).
- **Data:** repository impls with **fake data sources** (mapping DTO↔entity, cache/offline logic, exception→`Failure` conversion — Module 40).
- **Presentation:** view models via **state sequences** (loading → data/empty/error) with fake use cases (`blocTest` /listeners — Module 43).
- Each **isolated** with fakes at boundaries — the architecture's payoff (02_unit_and_mocking.md).

- **Widget + golden (middle):** widget-test the screen's **states + interactions** (faked VM); a **golden** for a stable design-system component/key screen (03_widget_and_golden_tests.md).

- **Integration (tip):** **one** integration test for the feature's **critical journey** (deterministic data, stable finders — 04_integration_and_e2e.md).

- **Coverage — meaningful, not vanity:**

- Measure with `flutter test --coverage` (LCOV) → view in CI / tools (Codecov). Use it to **find untested behavior + unhappy paths**, not to chase **100%** (which rewards testing trivia + gives false confidence).
- Prioritize coverage of **domain/data/presentation logic** and **error/edge paths**; don't count generated/trivial code. A pragmatic target (e.g., high coverage of *logic*) beats a blanket 100%.

- **CI gating** (Module 50):

- CI runs `flutter analyze` + `flutter test` (+ `--coverage`) on every PR; `integration_test` on emulators/device farm (perhaps on merge/nightly due to cost).
- Monorepo: `melos run test --since` for incremental, changed-package-only runs (Module 45).

- **Gate merges** on green tests (+ optionally a coverage threshold/no-regression); flaky tests quarantined, not tolerated.
- **Test organization:** mirror the source structure (`test/features/<name>/{domain,data,presentation}` , `integration_test/`); shared fakes/helpers in a `test/support` or a test package; consistent naming.
- **The payoff:** a **pyramid-shaped, layer-complete, CI-gated, meaningfully-covered** suite → confident refactoring + regression safety + fast feedback — realizing the testability the architecture band designed for (Module 47).

Memory Representation

The suite: unit tests (fakes + assertions) per layer, widget/golden tests (faked state + references), one integration test (real app + seeded data), a coverage report, and CI config. Organized to mirror the feature slice.

Compiler Behavior

Unit tests compile against interfaces + fakes; widget tests against real widgets + fakes; integration against the real app. Coverage instruments the code; CI compiles + runs everything.

Runtime Behavior

Unit (ms) → widget (test binding) → integration (device) run in ascending cost; CI runs unit+widget on every PR, integration on emulators (merge/nightly). Coverage collected during test runs. Merges blocked on failures.

Flutter Engine Behavior

Unit tests don't touch the engine; widget tests use the fake test binding; integration uses the real engine on-device — as in their respective files.

Dart VM Behavior

Fast unit base parallelizes in CI; incremental (`melos --since`) runs only changed packages' tests — keeping CI quick as the app scales (Module 45).

Examples

Test organization (mirrors the feature slice):

```
test/features/profile/
  domain/get_profile_test.dart      (use case + fake repo)
  data/profile_repository_test.dart (repo + fake sources; mapping/errors)
  presentation/profile_vm_test.dart (state sequences, fake use case)
  presentation/profile_view_test.dart (widget: states + interactions)
  presentation/profile_view_golden_test.dart (golden, stable UI)
integration_test/profile_journey_test.dart (one critical journey)
test/support/fakes.dart            (shared fakes/helpers)
```

CI (GitHub Actions sketch) – gate merges on green + coverage

- run: flutter analyze

- run: flutter test --coverage # unit + widget (+ LCOV)

- run: melos run test --since=origin/main # monorepo: changed packages only

- run: flutter test integration_test # on emulator (merge/nightly)

upload coverage; fail PR on test failure (and optionally coverage regression)

// Layer coverage at a glance (one test per layer, all fakes)

```
test('domain: use case maps repo Result', () async { /* fake repo */ });
```

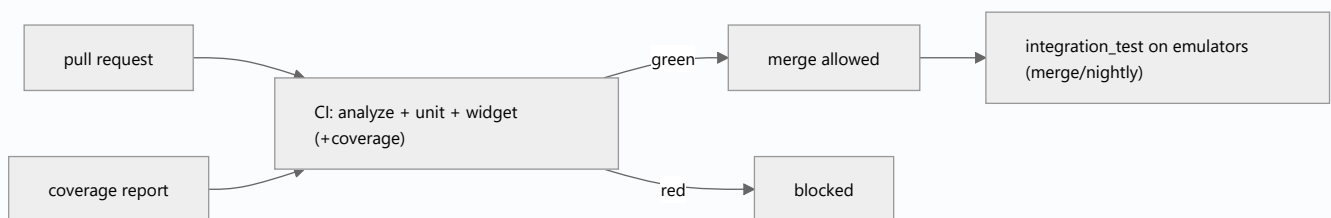
```
test('data: repo converts 404 -> NotFoundFailure', () async { /* fake sources */ });
```

```
blocTest('presentation: loading -> data', /* fake use case, assert sequence */);
```

```
testWidgets('widget: error state shows retry', (t) async { /* faked VM */ });
```

// integration_test: login -> profile (one critical journey)

Diagrams



Common Mistakes

Mistake	Why it's wrong	Fix
Testing only one layer (e.g., only UI)	Gaps + slow/imprecise	Unit-test each layer (domain/data/presentation)
Wrong shape (mostly E2E)	Slow/flaky CI	Pyramid: unit-heavy, thin E2E
Chasing 100% coverage	Rewards trivia, false confidence	Meaningful coverage (logic + unhappy paths)
Not gating merges on tests	Regressions slip in	CI gate on green tests
Running full E2E on every PR	Slow pipelines	E2E on merge/nightly; incremental unit/widget on PR
No shared fakes/organization	Duplication, hard to navigate	Mirror structure + shared test support
Tolerating flaky tests	Erodes trust in CI	Quarantine + fix fast
Not using incremental tests in a monorepo	Slow CI at scale	<code>melos run test --since</code>

Best Practices

- Build a **pyramid-shaped, layer-complete** suite: **unit-test each layer** (domain/data/presentation with fakes), **widget+golden** the UI, **one integration test** per critical journey.
- Measure **coverage meaningfully** (logic + unhappy paths; not 100% vanity; skip generated/trivial); use it to find gaps.
- **Gate merges in CI** (`flutter analyze` + `flutter test --coverage` on PRs; `integration_test` on emulators at merge/nightly; `melos --since` incremental in monorepos); **quarantine flaky** tests.
- **Organize** tests to mirror the source, with shared fakes/helpers; leverage the **testable architecture** so most coverage comes from fast unit tests.

Performance

The pyramid + incremental CI (`melos --since`) keep the suite fast: unit/widget on every PR (quick feedback), E2E rationed to emulators at merge/nightly. Meaningful coverage avoids wasted effort on trivia. Fast, gated tests = tight, safe iteration (Module 47).

Advantages / Disadvantages

- + Full-layer regression safety, confident refactoring, fast feedback (pyramid + incremental CI), meaningful coverage, merge gating.
- – Upfront + ongoing test-writing/maintenance, CI setup (device farm for E2E), coverage-target judgment, flake management.

Interview Questions

1. 🟢 **How do you test each architectural layer?** — Domain use cases with fake repos, data repos with fake sources, presentation view models via state sequences with fake use cases — all fast unit tests, plus widget/golden for UI and one E2E for the critical journey.
2. 🟢 **What shape should the suite be?** — A pyramid: many unit, some widget, few integration/E2E — keeping CI fast and reliable.
3. 🟡 **Why not chase 100% coverage?** — It rewards testing trivia and gives false confidence; aim for meaningful coverage of logic + unhappy paths, skipping generated/trivial code.
4. 🟡 **How do you gate merges with tests in CI?** — Run analyze + unit + widget (+coverage) on every PR (block on failure); run E2E on emulators at merge/nightly; incremental (`melos --since`) in monorepos.
5. 🟡 **Why run E2E on merge/nightly rather than every PR?** — E2E is slow/flaky; running it on every PR bloats the pipeline — run fast unit/widget on PRs, E2E less frequently.
6. 🔴 **How does testable architecture make this easier?** — Clean/MVVM (pure domain, view-model state, repository interfaces + fakes) pushes most logic into fast unit tests, so the pyramid's base is easy to build.
7. 🔴 **How do you keep CI fast + trustworthy as the app scales?** — Incremental changed-package test runs (`melos --since`), a unit-heavy pyramid, and aggressive flake quarantine/fix.

Senior Engineer Tips

- Build the base first: a fast unit test per layer (domain/data/presentation with fakes) delivers most of the safety cheaply; then add widget/golden and exactly one E2E per critical journey.
- Use coverage as a gap-finder, not a target — a green 100% suite that skips unhappy paths is worse than an 80% one that covers errors and edges.
- Gate PRs on fast tests, run E2E less frequently on emulators, use `melos --since` in monorepos, and quarantine flaky tests immediately — a slow or flaky CI is a suite people route around.

Architect Perspective

This capstone is where the handbook's testable architecture cashes out: a pyramid-shaped suite that unit-tests every layer (thanks to fakes at abstraction boundaries), guards UI behavior + visuals with widget/golden tests, validates critical journeys with a thin E2E tip, measures coverage meaningfully, and gates merges in fast incremental CI. The result is regression safety + confident refactoring + fast feedback at scale — the concrete enabler of everything from clean architecture to scalable applications and CI/CD (Module 40, Module 47, Module 50).

Summary

- Assemble a pyramid-shaped, layer-complete suite: unit-test domain/data/presentation (fakes), widget+golden the UI, one integration test per critical journey.
- Measure coverage meaningfully (logic + unhappy paths, not 100% vanity); gate merges in CI (analyze + unit + widget on PRs; E2E at merge/nightly; `melos --since` incremental).
- Mirror source structure + shared fakes; quarantine flakes; the payoff is confident refactoring + regression safety + fast feedback.

Revision Notes

- Per-layer unit: domain (fake repo), data (fake sources; mapping/errors), presentation (state sequences, fake use case) — isolated, fast. + widget/golden (faked VM, states/interactions/visuals) + 1 integration test (critical journey, deterministic).
- Coverage: `flutter test --coverage` (LCOV) → find gaps (logic + unhappy paths), not 100% vanity; skip generated/trivial.
- CI: analyze + unit + widget (+coverage) on PR (gate merge); `integration_test` on emulators at merge/nightly; `melos run test --since` incremental (monorepo); quarantine flakes; mirror source structure + shared fakes.

Practice Questions

1. What test type covers each architectural layer, and why?
2. Why measure coverage but not chase 100%?
3. How should tests be gated/run in CI (PR vs merge/nightly)?

Coding Questions

1. Write one unit test per layer (domain/data/presentation) using fakes.
2. Add a widget test (states + interaction) + a golden, and one integration test.
3. Write the CI config gating PRs on analyze + unit + widget (+coverage), E2E on merge.

Mini Project

Layered test suite + CI (capstone — Flutter): For a feature slice, write: unit tests per layer (use case + fake repo; repo + fake sources incl. error mapping; view model state sequences), widget tests (states + interaction) + one golden, and one integration test for the critical journey — organized to mirror the source with shared fakes. Add coverage (`--coverage`) used to find gaps and a CI config gating PRs on analyze + unit + widget (E2E on merge/nightly; `melos --since` if

monorepo). Acceptance: pyramid-shaped, every layer unit-tested with fakes; widget+golden UI; one critical-journey E2E; meaningful coverage (logic + unhappy paths, not 100%); CI gates merges (fast on PR, E2E less often); flakes quarantined; mirrors source structure.