



Master Study Notes: NPTEL

Software Engineering Course Notes

Table of Contents

- [Lecture 01: Introduction- I](#)
- [Lecture 02: Introduction- II](#)
- [Lecture 03: Introduction- III](#)
- [Lecture 04: Introduction- IV](#)
- [Lecture 05: Introduction- V](#)
- [Lecture 06: Life Cycle Model](#)
- [Lecture 07: Life Cycle Model \(Contd.\)](#)
- [Lecture 08: Waterfall Model](#)
- [Lecture 09: Waterfall Derivatives](#)
- [Lecture 10: Incremental Model](#)
- [Lecture 11: Evolutionary Model](#)
- [Lecture 12: Agile Model](#)
- [Lecture 13: Extreme Programming and Scrum](#)
- [Lecture 14: Scrum](#)
- [Lecture 15: Introduction to requirement specification](#)
- [Lecture 16: Requirement gathering and analysis](#)
- [Lecture 17: Functional requirements](#)
- [Lecture 18: Representation of complex programming logic](#)
- [Lecture 19: Design Fundamentals](#)
- [Lecture 20: Modular Design](#)
- [Lecture 21: Classification of Cohesion](#)
- [Lecture 22: Classification of Coupling](#)
- [Lecture 23: Introduction to structured analysis and structured design](#)
- [Lecture 24: Basics of Data Flow Diagrams \(DFD\)](#)
- [Lecture 25: Developing DFD Model](#)
- [Lecture 26: Examples of DFD Model development](#)
- [Lecture 27: DFD Model - More Examples](#)

- [Lecture 28: Essentials of Structure Chart](#)
 - [Lecture 29: Structure Chart Development](#)
 - [Lecture 30: Structured Design Examples](#)
 - [Lecture 31: Use Case Modelling](#)
 - [Lecture 32: Factoring Use Cases](#)
 - [Lecture 33: Overview of Class diagram](#)
 - [Lecture 34: Inheritance relationship](#)
 - [Lecture 35: Association relationship](#)
 - [Lecture 36: Aggregation/ Composition and dependency relations](#)
 - [Lecture 37: Interaction Modelling](#)
 - [Lecture 38: Development of Sequence diagrams](#)
 - [Lecture 39: State-Machine diagram](#)
 - [Lecture 40: An Object-Oriented design process](#)
 - [Lecture 41: Domain Analysis](#)
 - [Lecture 42: Examples of object-oriented design](#)
 - [Lecture 43: Basic concepts in Testing-I](#)
 - [Lecture 44: Basic concepts in Testing-II](#)
 - [Lecture 45: Basic concepts in Testing-III](#)
 - [Lecture 46: Unit testing strategies-I](#)
 - [Lecture 47: Unit testing strategies-II](#)
 - [Lecture 48: Equivalence Class Testing-I](#)
 - [Lecture 49: Equivalence Class Testing-II](#)
 - [Lecture 50: Special Value Testing](#)
 - [Lecture 51: Combinatorial Testing](#)
 - [Lecture 52: Decision Table Testing](#)
 - [Lecture 53: Cause Effect Graphing](#)
 - [Lecture 54: Pairwise Testing](#)
 - [Lecture 55: White box Testing](#)
 - [Lecture 56: Condition Testing](#)
 - [Lecture 57: MC/DC Coverage](#)
 - [Lecture 58: MC/DC Testing](#)
 - [Lecture 59: Path Testing](#)
 - [Lecture 60: Dataflow and Mutation Testing](#)
-

Lecture 01: Introduction- I

Part 1: அறிமுகம் மற்றும் அடிப்படை மன மாதிரிகள் (Introduction & Mental Models)

சாஃப்ட்வேர் இன்ஜினியரிங் (Software Engineering) துறைக்கான சரியான பொறியியல் அணுகுமுறை என்ன மற்றும் அது மரபுசார் அணுகுமுறையிலிருந்து எவ்வாறு வேறுபடுகிறது என்பதைப் புரிந்துகொள்வதே இதன் முதன்மை நோக்கமாகும்.

NOTE

> ****The Building Analogy (கட்டிட ஒப்புமை):**** > உங்கள் அடிப்படை உள் உணர்வுகளை (intuition) மட்டும் கொண்டு ஒரு சிறிய சுவரை உங்களால் சுலபமாகக் கட்டி முடிக்க முடியும். ஆனால், 30 முதல் 40 தளங்கள் கொண்ட ஒரு பெரிய கட்டிடத்தைக் கட்டச் சொன்னால், வெறும் உணர்வுகள் மட்டும் போதாது. முறையான பொறியியல் திட்டமிடல் இன்றி பெரிய கட்டிடத்தைக் கட்டத் தொடங்கினால், அந்தக் கட்டிடம் முழுவதும் தகர்ந்து போகலாம் அல்லது தரமற்றுப் போகலாம். > ****முக்கிய கருத்து:**** இதேபோலத்தான், பெரிய அளவிலான சாஃப்ட்வேரை அடிப்படை உணர்வை மட்டும் கொண்டு உருவாக்குவது, பொறியியல் இன்றி பெரிய கட்டிடத்தை உருவாக்குவது போன்ற ஒரு மோசமான தாக்கத்தை ஏற்படுத்திவிடும். அதனாலேயே முறையான சாஃப்ட்வேர் இன்ஜினியரிங் அணுகுமுறை தேவைப்படுகிறது.

Part 2: சாஃப்ட்வேர் நெருக்கடி (The Software Crisis)

சாஃப்ட்வேர் இன்ஜினியரிங் கொள்கைகள் என்பவை கடந்த கால அனுபவங்களிலிருந்து பெறப்பட்டவையாகும். முறையான பொறியியல் அணுகுமுறை இல்லாததால் "சாஃப்ட்வேர் நெருக்கடி" உருவானது.

சாஃப்ட்வேர் நெருக்கடி என்றால் என்ன? பயனாளர்களின் தேவையின் கடினத் தன்மை (complexity) அதிகரிக்கும் போது, அந்தத் தேவைகளை நிவர்த்தி செய்ய முடியாத நிலை ஏற்படுவதே **சாஃப்ட்வேர் நெருக்கடி (Software Crisis)** எனப்படும்.

IMPORTANT

> ****சாஃப்ட்வேர் நெருக்கடியின் அறிகுறிகள் மற்றும் புள்ளிவிவரங்கள் (Symptoms & Statistics):**** > ****காலதாமதம்:**** சாஃப்ட்வேரை காலம் தாழ்த்தி வெளியிடுவது வாடிக்கையாகிவிட்டது. > ****வளங்களின் வீணடிப்பு:**** வளங்கள் (Resources) உகந்த முறையில் உபயோகிக்கப்படுவதில்லை. > ****பட்ஜெட் மீறல்:**** சாஃப்ட்வேர் உருவாக்கத்தின் விலை வரம்பு மீறி உயர்கிறது. > ****வெற்றி விகிதம் (Success Rate):**** சாஃப்ட்வேர் செயல் திட்டங்களின் வெற்றி விகிதம் வெறும் ****28 சதவீதம்**** மட்டுமே. > ****ரத்து செய்யப்படுதல் (Cancellation):**** ****கால் பங்கு (25%)**** செயல்திட்டங்கள் வெளிவருவதே இல்லை; அவை முற்றிலுமாக ரத்து செய்யப்படுகின்றன. > ****தாமதம் &**

அதிக செலவு:** மீதமுள்ள பாதி செயல்திட்டங்கள் தாமதமாக வெளியிடப்படுகின்றன, மேலும் அதன் விலை பல மடங்கு உயர்கிறது.

Part 3: சாஃப்ட்வேர் vs ஹார்டுவேர் (Software vs. Hardware)

ஏன் வன்பொருளை (Hardware) விட மென்பொருள் (Software) அதிக அளவு தேர்ந்தெடுக்கப்படுகிறது என்பதற்கான ஆழமான காரணங்கள் விவாதிக்கப்பட்டுள்ளன.

ஒப்பீட்டு அளவுரு	ஹார்டுவேர் (Hardware)	சாஃப்ட்வேர் (Software)
விலை (Cost)	தற்போது ஹார்டுவேரின் விலை மலிவாக உள்ளது.	சாஃப்ட்வேர் உருவாக்கத்திற்கு அதிக செலவாகிறது.
உருவாக்க காலம் (Development Time)	தேவைகளை அறிவது முதல் கட்டுருவாக்கங்களை (prototypes) உருவாக்குவது வரை பல வருடங்கள் தேவைப்படுகின்றன.	ஹார்டுவேரை விட விரைவாக உருவாக்க முடியும்.
மாற்றங்களை கையாளுதல் (Adaptability to Change)	சிறிய மாற்றங்களைச் செய்ய வேண்டுமென்றாலும், வடிவத்தை மாற்றுவதற்கும் நீண்ட செயல்முறை மற்றும் வெகு நீண்ட வருடங்கள் தேவைப்படும்.	மாற்றங்களை சுலபமாக ஏற்றுக்கொள்ளும் தன்மை (malleability) கொண்டது.

The "Why": ஹார்டுவேர் மூலம் ஒரு தீர்வை உருவாக்குவது கடினமான மற்றும் மாற்ற முடியாத (inflexible) ஒன்றாகும். ஆனால் சாஃப்ட்வேரில் ஒவ்வொரு மாற்றத்திற்கும் ஏற்ப அதன் கட்டமைப்பில் (Software Structure) மாற்றங்களைச் செய்து கொண்டே இருக்க முடியும் என்பதே சாஃப்ட்வேரைத் தேர்ந்தெடுப்பதற்கு முக்கியக் காரணமாகும்.

Part 4: சாஃப்ட்வேர் உருவாக்கத்தின் சவால்கள் (Challenges in Software Development)

சாஃப்ட்வேர் உருவாக்குவது ஒரு சவால்கள் நிறைந்த கலையாகும். இதற்கான முக்கிய காரணங்கள்: 1. **கடினத்தன்மை மற்றும் பருமன்:** சாஃப்ட்வேர் உருவாக்குவது எளிதல்ல; அது பெரியதாகவும், அதிக ஆற்றலை உட்கொள்ளும் பருமனானதாகவும் இருக்கும். 2. **கையேடு**

முயற்சி (Manual Effort): இது முற்றிலுமான கையேடு (manual) முயற்சியைக் கோருகிறது. 3. **அறிவுப் பற்றாக்குறை:** புரோகிராமர்கள் மற்றும் உருவாக்குனர்கள் போதுமான ப்ரோக்ராம் உத்திகளில் அறிவுடனும் நிபுணத்துவத்துடனும் இருப்பதில்லை. 4. **மெதுவான உற்பத்தித்திறன் வளர்ச்சி (Slow Productivity Growth):** பிரச்சனைகளின் அளவும் கடினத்தன்மையும் அதிவேகமாக வளரும் அதே வேளையில், சாஃப்ட்வேர் இன்ஜினியரிங் ஆக்கத்திறன் (productivity) முன்னேற்றம் அதிதீவிரமாக இல்லை; இது வருடாவருடம் மிதமாகவும் மெதுவாகவுமே முன்னேறுகிறது.

Part 5: சாஃப்ட்வேர் - கலையா? கைவினையா? பொறியியலா? (Evolution of Software Engineering)

சாஃப்ட்வேர் உருவாக்கம் என்பது காலப்போக்கில் எப்படி ஒரு பொறியியல் துறையாக உருவெடுத்தது என்பதைப் புரிந்துகொள்ள, இரும்பு அல்லது காகிதம் உருவாக்கும் துறையின் தொழில்நுட்ப முன்னேற்றப் பாங்குடன் ஒப்பிடப்படுகிறது.

TIP

> ****பரிணாம வளர்ச்சியின் 3 படிநிலைகள் (The 3 Stages of Evolution):**** > > ****படிநிலை 1: கலை (Art)**** > ****உதாரணம்:**** ஆரம்பத்தில் காகிதம் செய்வது சீன மக்களுக்கு மட்டுமே தெரிந்த ஒரு கலையாக இருந்தது; அவர்களால் மட்டுமே நல்ல தரமான காகிதத்தை உருவாக்க முடிந்தது. மற்றவர்கள் உருவாக்கியவை தரம் குறைந்தவையாகவே இருந்தன. > ****ஒரு ஓவியரிடம் எப்படி அழகிய ஓவியம் வரைந்தீர்கள் என்று கேட்டால், "அது தானாக வந்தது (internal intuition)" என்று கூறுவார். இதுவே கலை.** > > ****படிநிலை 2: கைவினை (Craft)**** > ****தொழில்நுட்ப வளர்ச்சியுடன் கலையானது கைவினை வடிவமாக (Craft) மாறியது.** > ****செயல்முறை:**** இந்த வடிவத்தில் சில உத்திகளை (techniques) அடையாளம் காண முடிந்தது. ஆனால், வடிவமைப்பாளர்கள் அந்த உத்திகளைத் தங்களுக்குக் கீழ் வேலை செய்யும் தொழில் பழகுபவர்களுக்கு (apprentices) மட்டுமே பகிர்வார்கள். > ****இது சிறிய எண்ணிக்கையிலான மக்களிடம் மட்டுமே பகிரப்பட்ட ஒரு மறைக்கப்பட்ட ரகசியம் (hidden secret) ஆகும்.** > > ****படிநிலை 3: பொறியியல் (Engineering)**** > ****மெதுவாக இது பொறியியல் அணுகுமுறைக்கு மாறியது.** > ****செயல்முறை:**** நுணுக்கங்கள் அனைத்தும் அறிவியல் பூர்வமாக ஆய்வு செய்யப்பட்டன. பழைய அனுபவங்கள் கூர்ந்து ஆராயப்பட்டு, அவற்றுக்கு அறிவியல் அடிப்படைகள் (scientific foundations) கொடுக்கப்பட்டன. > ****விளைவு:**** இவை அனைத்தும் முறையான இலக்கிய நூல் வடிவமாக ஆவணப்படுத்தப்பட்டன (Documented). இதனால், வெளிப்படை அறிவுள்ள எவராலும் அதைப் படித்துப் பயில முடிந்தது.

சாஃப்ட்வேரின் பரிணாமம்: ப்ரோக்ராம் எழுதுவதும் இதே பாங்கையே (Pattern) பின்பற்றியது. 1960-களின் ஆரம்பத்தில் சிறந்த புரோகிராமர்கள் உபயோகப்படுத்திய நுட்பங்கள் முறையாக

ஆய்வு செய்யப்பட்டு, அறிவியல் அடிப்படைகள் புகுத்தப்பட்டு ஆவணங்களாக வெளியிடப்பட்டன. இதனாலேயே, மாற்றங்களுக்குத் தகுந்தாற்போல் அணுகக்கூடிய சாஃப்ட்வேர் இன்ஜினியரிங்கை எவரும் படிக்கவும் கற்றுக்கொள்ளவும் முடியும் என்ற நிலை உருவானது.

Lecture 02: Introduction- II

Part 1: The Transition from Exploratory to Engineering Approaches

In the early days of computing, software development relied entirely on the intuition of programmers because formal techniques did not exist [1]. This intuition-based methodology is formally known as the **Exploratory Style** or the **Build and Fix** approach [1, 2].

- **The Build and Fix Cycle:** The programmer writes code and then continuously fixes errors until the problem requirements appear to be met [2].
- **The Scaling Problem:** While the exploratory style is suitable for writing small programs, it fails catastrophically for large-scale, team-based software projects [2]. As program size increases, the **cost, effort, and time** required escalate at an exponential rate, eventually causing the exploratory method to break down completely [3, 4]. Software engineering techniques, in contrast, can reduce development time drastically (e.g., taking one-tenth of the time) without exponential cost scaling [3].

IMPORTANT

> **Why does the Exploratory Style fail for large projects?** > The root cause lies not in the machines, but in the limitations of the **human cognitive mechanism** [5]. Software engineering principles are explicitly designed to overcome these human limitations [4, 5].

Part 2: Human Cognitive Limitations

To understand why **uncontrolled complexity breaks down software development**, one must understand human memory [5].

- **Long-Term vs. Short-Term Memory:** Humans possess **Long-Term Memory** (which can store information for months or years) and **Short-Term Memory** (the processing center used for active computation and logic) [5-7]. Information in short-term memory decays very rapidly, typically lasting only around 10 seconds unless actively refreshed (e.g., repeating a grocery list) [7, 8].
- **Miller's Magic Number:** The capacity of human short-term memory is strictly limited. According to psychological principles (often referred to as George Miller's law), the human brain **can only hold and process about 7 items (or chunks)** at any given time [9, 10]. If a program's structure exposes the programmer to more than 7 interacting elements simultaneously, comprehension becomes incredibly difficult [9].

Part 3: Overcoming Cognitive Limits via Chunking and Abstraction

To handle software complexity and bypass the biological limit of 7 items, humans use a process called **Chunking** [11].

- **Chunking:** If discrete items have a logical connection, the brain groups them into a single "chunk" [8]. For example, random letters have no connection, but arranged correctly, they form a word (1 chunk). A binary string like `110010101001` is hard to memorize, but grouping it into 3-bit octal blocks makes it easily manageable [11, 12].

To build large systems without overwhelming short-term memory, software engineering relies heavily on two core principles to facilitate chunking: **Abstraction** and **Decomposition** [10] (Decomposition is expanded upon in Lec 03).

NOTE

> **Core Concept: Abstraction** > **Abstraction** is the process of focusing on the essential features of a system while intentionally ignoring irrelevant details to reduce complexity [10, 13].

- **Analogy - Building Architecture:** When creating the front elevation view of a multi-story building, you ignore the internal infrastructure and wall thickness. When creating the floor plan, you ignore the front elevation and focus purely on layout. By

ignoring irrelevant details for a specific context, you simplify the problem to a manageable level [10, 13, 14].

Lecture 03: Introduction- III

Part 1: Further Explorations of Abstraction and Decomposition

Software engineering introduces systematic techniques to deal with the rapidly increasing complexity of software [15].

- **Abstraction Analogies:**
 - *The Country Map:* If tasked with understanding a country's market, physically visiting every location would take up to 100 years and still yield an incomplete picture [16]. Instead, you use map abstractions. A roadmap shows basic transit models; a physical map shows mountains, rivers, and oceans. Each abstraction focuses on one specific aspect, ignoring the rest [17].
 - *Biological Classification:* Understanding the biological design of all living organisms is impossible in a single lifetime. Instead, we use hierarchical abstractions and models (like a book's table of contents) to understand life in stepwise levels [18].
- **Decomposition:** This is the process of breaking a large, complex problem into smaller, highly manageable, independent parts [15, 19].
 - *Analogy:* It is difficult to break a bundled set of sticks, but breaking them individually is easy. Similarly, a continuous, unstructured book is impossible to read; breaking it into structured chapters makes it readable [15, 19].

Part 2: Software Project Characteristics

How do software projects differ from general research or exploratory projects? *

Predictability: Exploratory research (like finding a cure for cancer) has uncertain outcomes; you may work for years without success [20]. A software project, however, is

deterministic. It is a well-defined task with known challenges and repeatable routines that *can* be successfully completed [20, 21].

 TIP

> **Types of Software Projects** > 1. **Product Development Projects (Horizontal Market):** Software built to be sold to a large number of different organizations in the open market [22]. > 2. **Custom / Service Projects (Vertical Market):** Software developed specifically for a single client (e.g., a specific bank) or modified from existing software to meet niche requirements [22].

Lecture 04: Introduction- IV

Part 1: Industry Context and System Engineering

- **The Services Boom:** In the Indian IT industry (backed by NASSCOM data), service-oriented software projects see massive growth [23, 24].
- **Software Reuse:** Modern software is rarely built entirely from scratch. A significant portion of development involves reusing existing codebases and customizing them, adding new functionalities step-by-step for the client [24].
- **Computer System Engineering:** Software rarely exists in a vacuum; it is usually part of a larger hardware-software system [25]. The engineering process involves a high-level understanding of the system, developing the software, testing it via simulation, and finally integrating it with the hardware to test as a complete system [25].

Part 2: The Evolution of Programming Practices

Programming techniques evolved from simple, unstructured methods to highly refined engineering practices to manage growing complexity [26].

- **1950s - Assembly Language:** Programming was done in assembly language, heavily relying on the programmer's intuition. Programs were small (around 1,000 lines), took a long time to finish, and were highly error-prone [26, 27].

- **1960s - High-Level Languages (HLLs):** Languages like FORTRAN drastically increased productivity. One line of HLL code could handle the equivalent of several assembly instructions. Furthermore, programmers no longer had to manually manage internal hardware structures (like registers) [26, 27].

IMPORTANT

> **The Problem with Control Flow & GOTO** > As program size increased, developers realized that the execution sequence (**Control Structure**) became too difficult to track [27]. Early assembly and HLL code relied heavily on `GOTO` statements, creating complex, tangled paths (Spaghetti code) that were impossible to trace, understand, or debug [28, 29]. This led to Dijkstra's famous stance that the `GOTO` statement is harmful [28].

- **Structured Programming:** To fix the chaotic control flow, structured programming was introduced. It dictates that code should be written without `GOTO` statements, relying purely on structured paths. This makes the code simpler to read, understand, and maintain, drastically reducing the number of bugs [29, 30].

Lecture 05: Introduction- V

Part 1: Deep Dive into the Evolution of Paradigms

This lecture synthesizes the historical progression of software design techniques, highlighting exactly *why* certain paradigm shifts occurred.

1. High-Level Languages (1960s)

Languages like FORTRAN, ALGOL, and COBOL replaced Assembly language for three fundamental productivity reasons [31, 32]:

1. **Hardware Abstraction:** HLLs abstract away machine architecture and registers, allowing programmers to write code using variables that map to real-world concepts [31, 32].
2. **Instruction Density:** Every single high-level construct equates to 3 to 4 assembly instructions, meaning less typing and fewer lines to maintain [32].
3. **Human Readability:** HLLs are significantly closer to natural human language than assembly, making them inherently easier to write and read [32].

2. Structured Programming & Flow-Charting

Despite HLLs, the 1960s still relied on the "exploratory" style. As program sizes grew to thousands of lines, intuition failed, and bugs proliferated [33]. * **Flow-Charting:** Introduced to design the **Program Control Structure** (the sequence of executable statements) *before* coding [34]. Without modeling logic first, control structures become poor, making tracing inputs and outputs nearly impossible, even if one spends years analyzing the code [35, 36]. Flow-charting limits execution paths, speeding up development and debugging [36]. * **The 3 Core Constructs:** Dijkstra's 1969 letter ("GOTO considered harmful") proved that any programming logic can be expressed using only three structures: 1. **Sequence** 2. **Selection** (Conditionals) 3. **Iteration** (Loops) [37, 38]. *Note: GOTO is only deemed acceptable for practical edge cases, like premature loop exits or exception handling [39].*

Part 2: The Shift to Data and Objects (1970s - 1980s)

NOTE

> **Why Control Structure wasn't enough:** > Structured programming solved control flow, but developing large programs remained too slow and expensive. In the 1970s, the industry realized that **Data Structures** needed equal attention [40, 41].

- **Data Structure Oriented Design (1970s):** Methodologies like **Jackson's Structured Programming (JSP)** and the **Warnier-Orr** method became popular. The core objective was to design the data structure first (using diagrammatic notations for sequence, selection, and iteration) and then derive the program structure directly from the data structure [41, 42].
- **Data Flow Design (DFD):** Focuses heavily on how data flows from input, through processing stations, to output [43].
 - *Analogy - Car Assembly Line:* An engine and chassis enter a processing station and are joined. They flow to the next station to receive doors, then to a parallel line for wheels, finally exiting as a painted, ready-to-ship car [44, 45]. DFD uses simple notations to model these functions and data flows, taking less than an hour to learn [44].

3. Object-Oriented Design (1980s)

In the 1980s, Object-Oriented (OO) techniques emerged. * **Why OO?** It appealed to developers because it represents natural, real-world objects within the problem space [45]. * **Architectural Trade-offs & Benefits:** * OO leads to superior modular design because objects act as excellent modules that enforce **Data Hiding** and **Data Abstraction** [46]. * It simplifies the design of massive, complex problems [46]. * It strongly facilitates **Code Reuse**, which directly reduces development time and cost [46]. * It produces code that is fundamentally less buggy and far easier to maintain compared to earlier paradigms [46, 47].

Lecture 06: Life Cycle Model

Part 1: Evolution from Exploratory to Systematic Development

The software engineering paradigm has evolved significantly from the 1950s and 1960s, moving away from exploratory "build and fix" methods toward well-defined **Life Cycle Models** [1]. In earlier eras, bugs were addressed only after the entire program was written, leading to a tedious one-by-one removal process [1], [2]. Modern methodologies emphasize defect prevention and early detection [2]. Today, software development incorporates multiple phases, and reviews are conducted at the end of each phase so that errors are caught and corrected immediately, rather than waiting for the testing phase [2], [3]. This systematic approach significantly reduces overall development cost and time [3].

Part 2: The Necessity of a Life Cycle Model (LCM)

A **Life Cycle Model** breaks down the software creation process into distinct, manageable phases and establishes the exact sequence in which these phases should be executed [4], [5].

IMPORTANT

> Without a documented Life Cycle Model, development becomes chaotic. Team members may lack a unified understanding of who is doing what and when, leading to project failure [6].

Having a documented LCM provides a roadmap that helps identify missing, redundant, or conflicting activities [7]. It also grants the project manager increased visibility, enabling them to accurately estimate project completion times and required resources [8]. Furthermore, producing good documentation throughout the lifecycle phases makes future software maintenance significantly easier [8].

Part 3: Entry and Exit Criteria & Milestones

Each phase in a Life Cycle Model operates on strict boundary conditions: * **Entry Criteria:** The prerequisites that must be fulfilled before a phase can begin [9]. For instance, before starting the requirement specification phase, all requirements must be gathered, analyzed, documented, and reviewed by both team members and clients [10]. * **Exit Criteria:** The conditions that must be met to declare a phase complete [9].

NOTE

> When the exit criteria of a phase are successfully met, a **Milestone** is achieved [11]. Milestones are crucial for project managers to track progress, plan future tasks, and monitor the project's health [11].

Without these objective milestones, managers suffer from the **"99% Complete" Syndrome** [12]. This syndrome occurs when developers rely purely on overconfident intuition, repeatedly claiming the project is "99% done," while actual completion continuously drags on and drains manager patience [13], [12].

Lecture 07: Life Cycle Model (Contd.)

Part 1: The Classical Waterfall Model Overview

The **Classical Waterfall Model** is the most intuitive and traditional LCM [14]. It breaks the lifecycle into sequential phases: **Feasibility Study, Requirement Analysis and Specification, Design, Coding and Unit Testing, Integration and System Testing, and Maintenance** [15], [14]. If we analyze the effort distribution across the entire lifecycle of a product, the **Maintenance phase** consumes the maximum effort (often up to 60%) [16],

[17]. Within the development phases themselves, the **Testing phase** requires the most effort [16].

Part 2: Feasibility Study Phase

The first phase of the Waterfall model is the **Feasibility Study**, which determines if the project is worth undertaking. A project must be evaluated across three dimensions: 1. **Economic Feasibility (Cost-Benefit Analysis)**: Assessing if the financial benefits of the software outweigh the costs of development and operation [18], [19]. 2. **Technical Feasibility**: Determining if the organization possesses the required technical skills and resources to build the solution [20], [21]. 3. **Temporal/Time Feasibility**: Ensuring the project can be completed within the client's required timeframe [22], [23].

TIP

> ****Real-World Applicability:**** A case study of a mining company requiring a special Provident Fund system illustrates this process [24]-[25]. The project manager must visit the site, understand constraints, identify input data (daily/weekly contributions), and propose multiple architectural solutions (e.g., localized databases vs. central headquarters database) [26]-[27]. Each solution undergoes a cost-benefit analysis before deciding to proceed or abandon the project [27]-[28].

Part 3: The Business Case Document

The output of the feasibility study is a **Business Case**, which is submitted to upper management [29]. A good business case includes: * **Executive Summary**: A high-level overview for management [30]. * **Context and Business Opportunities**: Why the project is needed and the expected benefits [30]. * **Costs**: Expenses for development, deployment, operations, and training [30], [31]. * **Benefits**: Tangible (revenue generation, cost reduction) and intangible (faster processing, unquantifiable user benefits) [19], [29], [31]. * **Risks**: Identification of potential pitfalls, such as cost overruns, schedule delays, or lack of user adoption, along with mitigation strategies [31], [32].

Lecture 08: Waterfall Model

Part 1: Requirement Analysis and Specification

Following feasibility, the **Requirement Analysis** phase aims to gather and deeply understand user requirements [33]. During gathering, data is collected via interviews and discussions [34]. Analysts must resolve three primary anomalies during this stage: 1. **Contradictions:** When one requirement conflicts with another [35]. 2. **Ambiguities:** When requirements are unclear or vague [35], [34]. 3. **Incompleteness:** When necessary requirements are entirely missing [35], [34]. Once resolved, these requirements are formalized into the **Software Requirement Specification (SRS)** document [34].

Part 2: Design, Coding, and Testing

- **Design:** The SRS serves as the foundation for design [36]. Traditional approaches transform Data Flow Diagrams (DFDs) into structure charts, while Object-Oriented Design transforms object structures into detailed architectures [36]. Object-Oriented design is highly popular as it requires less time and effort and is easier to maintain [36], [37].
- **Coding:** Design is converted into source code, which is then documented [37].
- **Testing: Integration Testing** is performed to ensure that independently developed modules work together seamlessly without errors and fulfill all customer requirements [37], [17].

Part 3: Maintenance Phase

Maintenance begins after the software is deployed and continues for a long duration, taking up 60% of the lifecycle effort [17]. It is categorized into three types: 1. **Corrective Maintenance:** Fixing bugs discovered by users [38]. 2. **Perfective Maintenance:** Enhancing the software by adding new features that were not initially identified [38]. 3. **Adaptive Maintenance:** Modifying the software to run in new environments or hardware [38].

Part 4: Phase Decay and the Cost of Late Defect Detection

A massive drawback of the **Classical Waterfall Model** is the assumption that no errors are made in any phase [39]. If an error occurs during requirement gathering, it propagates to the design and coding phases [39], [40].

IMPORTANT

> The cost of fixing an error increases exponentially the later it is discovered. Fixing a requirement flaw during the coding or testing phase requires backward traversal to modify the SRS, the design, and the code, inflating costs massively compared to catching it in the initial phase [40], [41].

This critical flaw necessitated the evolution of the **Iterative Waterfall Model**, which incorporates feedback paths to allow backward corrections [42].

Lecture 09: Waterfall Derivatives

Part 1: Iterative Waterfall Model

The **Iterative Waterfall Model** allows phases to flow backward to correct errors discovered downstream [42]. Despite its improvements, it shares a core flaw with the classical model: it demands that all requirements be fully gathered and frozen before development begins [43], [44]. * **Drawbacks:** Customers cannot visualize the software beforehand, leading to missing, contradictory, or incorrect requirements [45], [46]. Once development starts, making changes forces developers to redo massive amounts of work (re-architecting and rewriting code) [45]. Furthermore, actual integration happens late, giving a false sense of progress and keeping the customer excluded until the very end [46], [47]. * **Applicability:** Best suited for well-understood domains, like accounting software, where requirements are static and the technical team is highly experienced [48], [47].

Part 2: The V-Model (Verification and Validation)

To address testing inefficiencies in the Waterfall model, the **V-Model** was introduced. It ensures that testing is spread across all development phases rather than left to the end [49]. * In the V-Model, testing plans are created parallel to development phases. For example, the system test plan is created during the requirement phase to prove that all gathered requirements are met [49]. * **Drawback:** Like Waterfall, phases cannot overlap, and it is strictly suited only for projects where requirements can be predicted completely upfront [50].

Part 3: Prototyping Model

For projects where technical risks are high or requirements are unclear, the **Prototyping Model** is used [50], [51]. * A "dummy" or preliminary version of the software (the prototype) is built and shown to the customer before actual development begins [50], [51]. * **Advantages:** Helps customers visualize the final product, clarifies requirements, reduces the need for heavy initial documentation, and significantly lowers maintenance costs by improving overall quality [51], [52].

Lecture 10: Incremental Model

Part 1: Shortcomings of Prototyping and Waterfall

While Prototyping solves many UI and requirement clarity issues, the initial prototype is usually discarded, adding an upfront cost [53], [54]. On the other hand, Waterfall models fail in modern environments because nearly 40% of requirements change *after* development has started [55]. Rigidly freezing requirements and building plans around them leads to massive cost overruns when changes inevitably occur [56], [57]. Fred Brooks noted that building systems sequentially with the intent to learn from failures is a practical reality [57].

Part 2: Incremental Model Architecture

The **Incremental Model** solves the rigidity of Waterfall by dividing the total requirements into smaller, manageable sub-modules or "features" [57], [58].

TIP

> **How it works:** Instead of building the whole system at once, developers take a small portion of the requirements, design it, code it, and deploy it to the customer. This deployed version is called an **Increment** [57], [58].

- **Feedback Loops:** Once an increment is deployed, the customer uses it and provides feedback, which is then used to refine future increments [58].
- **Execution:** The project iterates through design, build, and deploy phases for each increment until the full system is complete [58], [59].

Part 3: Prioritizing Increments (Risk and Value Formula)

To determine which increment to build first, project managers evaluate features based on a straightforward risk/value framework: 1. **Value:** How much value the increment provides to the customer [59]. 2. **Cost:** The effort and cost required to develop the increment [59]. Both metrics are commonly scored on a scale from 1 to 10 [60]. Features with the highest customer value and the lowest development risk/cost are prioritized for the earliest increments [59], [60].

Part 4: Model Comparison Table

Feature	Classical Waterfall	Iterative Waterfall	Incremental Model
Requirements	Must be frozen upfront.	Must be frozen upfront.	Flexible; broken into smaller feature sets.

Feature	Classical Waterfall	Iterative Waterfall	Incremental Model
Error Correction	Not possible; assumes no errors.	Possible via feedback loops to previous phases.	Constant correction via user feedback on increments.
Customer Involvement	Only at requirements and final deployment.	Only at requirements and final deployment.	Continuous; evaluates every deployed increment.
Delivery	Single massive deployment at the very end.	Single massive deployment at the very end.	Delivered in successive, usable versions.
Adaptability to Change	Extremely poor (high cost impact).	Poor (changes force massive rework).	Excellent (built to handle changing needs).

Lecture 11: Evolutionary Model

Part 1: The Need for the Evolutionary Model

In real-world software projects, requirements continuously change during the development phase [1]. Users often cannot visualize the software until it is built, leading to missing or ambiguous requirements [1, 2]. Researcher Caper Jones found that in a study of 8,000 projects, about 40% of requirements changed over time [2]. Traditional methodologies, like the **Waterfall Model**, lack the flexibility to handle these changing requirements effectively [2]. While the **Incremental Model** delivers software in parts, it still relies on the assumption that all requirements are fully known upfront [2, 3]. The **Evolutionary Model** was developed specifically to overcome the need to know all requirements in advance [3].

Part 2: Mechanics of the Evolutionary Model

The core philosophy of the **Evolutionary Model** is to "plan a little, design a little, code a little" [3]. * **Iterative Process:** Development happens in cycles called **iterations**, where essential features are developed, refined, and deployed incrementally [4]. * **Mini-Waterfalls:** Each iteration functions like a miniature Waterfall model involving coding, testing, and integration before being deployed to the user [5]. * **Timeboxing:** Each iteration typically lasts between 2 to 6 weeks, and a complete software project might take 10 to 15 iterations [5].

IMPORTANT

> A major advantage of the Evolutionary Model is that users can test the software and provide feedback early, allowing developers to course-correct before delivering the final product [6, 7].

Part 3: Advantages and Disadvantages

The Evolutionary Model embraces user feedback, allowing features that users dislike to be rejected and replaced with new ones [8].

Advantages: * Users get an actual feel for the software early on [7]. * Errors and defects are identified and reported during the continuous usage of core modules [7, 9]. * The final software is highly reliable and exactly matches user needs [7, 9]. * Easily accommodates changing requirements without long-term rigid plans [9].

Disadvantages: * The process is unpredictable and lacks a long-term vision [10, 11]. * It is difficult to estimate the total cost, duration, and required human resources [11]. * Continuous changes can degrade the software's structural integrity, potentially leading to an endless project loop [11, 12].

Part 4: Unified Process (UP)

The **Unified Process (UP)** is a popular iterative and incremental framework designed by Jacobson, Booch, and Rumbaugh, heavily utilized for object-oriented software development [13]. UP development occurs in four major phases: **Inception, Elaboration, Construction, and Transition** [14]. * **Inception:** Focuses on user communication, understanding features, and project planning [15]. * **Elaboration:** Focuses on creating models for the features [15]. * **Construction:** The software is built and executed [15]. * **Transition:** The software is deployed to the user's environment [15].

Lecture 12: Agile Model

Part 1: The Spiral Model (Risk Handling)

Before Agile, the **Spiral Model** was introduced by Boehm, combining features of the Waterfall, Incremental, and Evolutionary models (often called a meta-model) [16, 17]. The most distinguishing characteristic of the Spiral Model is its focus on **Risk Handling** [18, 19].

Each loop in the spiral represents a phase, divided into four quadrants: 1. **Quadrant 1 (Objectives & Alternatives)**: Identifies objectives and evaluates alternative solutions to resolve issues [19]. 2. **Quadrant 2 (Risk Analysis & Prototyping)**: Detailed analysis is conducted, and a prototype is built specifically to identify and resolve risks [17, 19]. 3. **Quadrant 3 (Development)**: Actual development and verification take place after risks are mitigated [17]. 4. **Quadrant 4 (Review)**: The phase is reviewed, though it does not always yield a deployable software increment to the user [17].

Part 2: Introduction to the Agile Model

Agile signifies rapid development [20]. It aims to eliminate time-wasting activities and unnecessary tasks [20]. For instance, in the traditional Waterfall model, up to 50% of the effort is spent on creating documentation that is rarely used [20, 21]. The Agile Model drastically reduces documentation, focusing instead on neatly incorporating user feedback [21]. Increments in Agile are very short, typically lasting 1 to 4 weeks [21].

 TIP

> Agile is actually an umbrella term that encompasses several methodologies, including Extreme Programming (XP), Scrum, Unified Process, Crystal, DSDM, and Lean [22].

Part 3: The Agile Manifesto & Principles

Published in 2000, the **Agile Manifesto** outlines key development values [21, 23]: * **Individuals and interactions** over processes and tools [23]. * **Working software (code)** over comprehensive documentation [23]. * **Customer collaboration** over contract negotiation [22, 23].

Core Agile Principles & Practices: * **User Stories:** Requirements are gathered informally using the user's own words rather than formal use cases [22, 24]. * **Spikes:** Exploratory prototypes used to evaluate potential solutions and alternatives [25]. * **Face-to-Face Communication:** Considered the most efficient communication method; facilitated by sharing office space (co-location) and using whiteboards [26, 27]. * **Minimal Documentation:** Only essential, precise, and highly detailed reports are generated [28].

Part 4: Agile vs. Traditional (Waterfall/Predictive) Models

Feature	Traditional/Waterfall Model	Agile Model
Planning	Long-term, rigid planning with upfront requirement gathering [29, 30].	Short-term planning with continuous requirement evolution [30].
Documentation	Heavy focus on extensive documentation and phase-end reports [30, 31].	Minimal documentation; relies on face-to-face communication [28, 32].
Delivery	Single, final deployment at the end of the lifecycle [33].	Incremental delivery in short cycles (1-4 weeks) [21].
Customer Involvement	Customers are isolated during the development phase [34].	Customers are integrated as part of the team for continuous feedback [23, 27].

NOTE

> ****Why Agile over Waterfall?**** Modern software projects face rapid changes, shorter timelines, and demand high customer satisfaction [35]. Agile's ability to adapt to

changes incrementally makes it far superior to Waterfall for such dynamic environments [35].

Lecture 13: Extreme Programming and Scrum

Part 1: Extreme Programming (XP)

Extreme Programming (XP) was introduced by Kent Beck in 1999 [36]. The core philosophy of XP is to take proven, effective software engineering practices and implement them at an "extreme" level [36].

Key XP Practices:

- * **Pair Programming:** Code review is highly effective, so XP takes it to the extreme by having two programmers write code together at one computer [37-39]. They switch roles every half hour [38].
- * **Test-Driven Development (TDD):** Testing ensures reliability. XP enforces TDD, where test cases are written *before* the code [40, 41]. Code is continuously refactored until it passes the tests [40, 41].
- * **Continuous Integration:** To avoid late integration issues, code is integrated and tested multiple times a day [42, 43].
- * **Refactoring:** Code is continuously improved and modified to enhance performance and structure without changing its behavior [42, 44].
- * **Simple Design:** Focus strictly on present needs rather than future complexities (Keep It Simple) [45, 46].
- * **Collective Ownership:** Any programmer is allowed to change another programmer's code to fix issues [39, 43].

XP Values: XP is driven by four core values: **Communication** (team and user), **Simplicity**, **Feedback** (continuous user input), and **Courage** (the willingness to discard bad code and rewrite it) [46, 47].

Part 2: Introduction to Scrum

Scrum is another highly popular Agile model [48]. Its most defining characteristic is the reliance on a **Self-Organizing Team**, where members autonomously decide who handles specific tasks without a head programmer dictating assignments [48].

- **Sprint:** The fundamental iteration in Scrum, typically lasting one month [48, 49].

- **Product Backlog:** A continuous repository where all software requirements (user stories) are stored [50].
 - **Daily Scrum:** A daily meeting where developers gather for 10-15 minutes to discuss progress and resolve immediate problems [50].
-
-

Lecture 14: Scrum

Part 1: Scrum Artifacts

Scrum manages requirements and tracks progress using specific artifacts. * **Product Backlog:** A dynamic list, usually maintained in a spreadsheet, containing all the User Stories (informal requirements) for the product [51-53]. It is continuously updated and prioritized [53, 54]. * **Sprint Backlog:** A subset of the Product Backlog selected for the current Sprint [53]. Once a Sprint begins, the Sprint Backlog is locked, and no external changes are allowed [55, 56]. * **Burndown Charts:** Used to track progress visually. * *Sprint Burndown Chart:* Tracks remaining hours of work within the current Sprint [57]. * *Release Burndown Chart:* Tracks progress toward the next release across multiple Sprints [58]. * *Product Burndown Chart:* Indicates the remaining work for the entire project [59].

Part 2: Scrum Ceremonies

Scrum defines specific timeboxed meetings (ceremonies) to manage the workflow: * **Sprint Planning:** A meeting where User Stories are selected from the Product Backlog to form the Sprint Backlog [60, 61]. The team agrees on realistic goals they can achieve in a month [62]. * **Daily Scrum:** A daily 15-minute stand-up meeting held in the morning [63, 64]. Members review what they did yesterday, what they plan to do today, and highlight any impediments [64, 65]. It is an information-gathering session, not a problem-solving meeting [64, 66]. * **Sprint Review:** Held at the end of the Sprint to review the developed increment [66]. It is an informal meeting (approx. 2 hours) where the team demonstrates the working software to stakeholders and the Product Owner [66].

Part 3: Scrum Roles

A Scrum team consists of three primary roles: 1. **Product Owner:** Represents the customer and the organization's interests [67, 68]. They manage the Product Backlog, prioritize features, determine release dates, and ultimately accept or reject the final increment [69]. 2. **Scrum Master:** Acts as the project manager and management representative [68, 69]. They remove impediments (e.g., network issues, hardware failures), shield the team from external interference, and facilitate communication [69, 70]. 3. **Development Team:** A cross-functional group of 5 to 10 members with expertise in coding, design, testing, etc [67, 70]. They are self-organizing and choose their own tasks [70].

Part 4: XP vs. Scrum

NOTE

> While both are Agile frameworks, they focus on different aspects of development. XP provides strict engineering guidelines, while Scrum provides a project management framework.

Feature	Extreme Programming (XP)	Scrum
Core Focus	Technical engineering practices (Pair programming, TDD) [37, 40].	Project management, team organization, and delivery [48, 68].
Iteration Length	Very short (days to a couple of weeks) [39, 45].	Timeboxed to 1 month (Sprint) [48, 51].
Requirement Changes	Highly adaptable to changes even late in the cycle [71].	Locked during a Sprint; no external influence allowed [56].
Roles	Customer is on-site and part of the team [43].	Distinct roles: Product Owner, Scrum Master, Team [72, 73].

Note: For extremely large projects, Scrum can be scaled using "Scrum of Scrums" or "Meta Scrum," where Scrum Masters and Product Owners from multiple teams coordinate overall features [59, 74].

Lecture 15: Introduction to requirement specification

Part 1: The Requirement Engineering Process

Requirement analysis and specification is one of the most critical and complex phases in the software development lifecycle [1]. The core objective is to understand the exact needs of the users and document them accurately to prevent costly failures later in the project [2].

IMPORTANT

> **The "Why" Behind Rigorous SRS:** Fixing a requirement error during the later phases of development (like design, coding, or testing) increases the cost exponentially [3], [4]. A misidentified requirement leads to cascading rework: modifying the design document, rewriting the code, and repeating the testing processes [4]. A high-quality SRS document minimizes these exponentially growing costs [5].

The requirement engineering process consists of three main iterative activities [6], [7]:

1. **Requirement Gathering (Elicitation):** Collecting requirements directly from the users' minds, understanding their needs, and mapping the problem scope [8], [9].

2. **Requirement Analysis:** Examining the gathered requirements to identify and remove problems, inconsistencies, and anomalies [8].

3. **Requirement Specification:** Documenting the finalized, problem-free requirements into the formal **Software Requirements Specification (SRS)** document [8], [9].

NOTE

> If an organization is building a generalized software product without a specific client, the marketing or sales team acts as the proxy "customer" to provide the requirements based on market needs [10], [11].

Part 2: Requirement Gathering Techniques

Analysts must employ structured techniques to extract accurate requirements from users, as users often do not fully understand the technical scope of their own problems [12], [13]. The analyst requires high creativity, imagination, and communication skills to suggest valuable features [12], [13].

The five primary techniques for requirement gathering are [14]-[15]:

1. **Studying Existing Documents:** If an existing manual system or an older software version exists, analyzing it provides a baseline understanding of operations [16].
2. **Interviews:** Conducting direct interviews with stakeholders, end-users, and clients to capture their expectations and daily operational needs [17], [18].
3. **Task Analysis:** Identifying the specific tasks that the system needs to perform based on the interview data, and formulating proper procedures for each task [17].
4. **Scenario Analysis:** Breaking down each task into various possible scenarios. For instance, a "leave request" task might have different scenarios like medical leave or semester withdrawal [17], [19].
5. **Form Analysis:** Analyzing the input forms submitted by users and the output forms generated by the current system to understand the exact data elements required [15], [20].

Lecture 16: Requirement gathering and analysis

Part 1: Requirement Analysis and Anomaly Resolution

The primary goal of requirement analysis is to meticulously examine the gathered requirements to identify and eliminate inherent problems [21]. When requirements are initially gathered, they typically contain three major types of anomalies [22]:

- **Ambiguity:** Requirements that are vague, lack precision, or can be interpreted in multiple ways [22], [23]. For example, stating "The same control field for all customers" is ambiguous—does it mean a single shared control field, or identical but separate fields for each customer? [23], [24].
- **Contradiction (Inconsistency):** Different requirements that logically clash with one another [24]. For example, one requirement states "turn on the heater when

temperature exceeds 100 degrees," while another states "turn off the heater when temperature exceeds 100 degrees" [25].

- **Incompleteness:** Missing requirements or unhandled edge cases that the customer forgot to specify [25], [26]. For example, defining system behavior when the temperature is over 100 degrees, but failing to specify what happens when it drops below 90 degrees [26].

Resolving these anomalies requires returning to the customer to clarify intents and ensure the system solves the actual business problem [27], [28].

Part 2: Characteristics of a Good SRS Document

Once anomalies are stripped away, the requirements are formally documented into the SRS [29]. The SRS acts as a legally binding contract between the client and developer, a reference for developers to write code, a guide for project managers to estimate costs/schedules, and a baseline for testers to write test cases [5], [30], [31].

TIP

> **Black-Box Specification:** The SRS must act as a **Black-Box** specification. It should specify *what* the system must do (input data and expected output behavior) rather than *how* the system will technically achieve it (internal logic) [32], [33].

A high-quality SRS document must possess the following properties [34]-[35]:

- * **Concise yet Exhaustive:** Minimal verbosity but omitting no critical details [34].
- * **Implementation Independent:** Should avoid specifying internal design or algorithms [34].
- * **Easily Modifiable:** Structured so that changing one functionality does not require hunting through the entire document [36].
- * **Consistent & Complete:** Free of contradictions and missing data [36].
- * **Traceable:** Requirements must be numbered or tagged with IDs so that design elements, code modules, and bugs can be explicitly traced back to the original requirement [37], [38].
- * **Testable:** Every requirement must be empirically verifiable. Subjective terms like "user-friendly" are impossible to test objectively [38], [39].

Lecture 17: Functional requirements

Part 1: Functional vs. Non-Functional Requirements

The SRS distinguishes between functional requirements, non-functional requirements, and constraints [40].

1. Functional Requirements: These represent the core meaningful work or facilities the software provides to the user [41]. Each functional requirement takes an input, processes it, and produces an output [42], [43]. * *Documentation structure:* Each functional requirement must detail its overview, input data, processing steps (validations/algorithms), expected output, and system behavior under abnormal/invalid input scenarios [44], [45].

2. Non-Functional Requirements (NFRs): These are system-wide characteristics and quality attributes that cannot be expressed as a single function [45], [46]. * *Examples:* Operating system compatibility (e.g., must run on UNIX and Windows), maintainability, security, reliability, and performance metrics (e.g., response time under 1 second for 90% of requests) [46], [47]. * *Rule:* All NFRs must be measurable and testable [48].

3. Constraints and Interfaces: * **Constraints:** Limitations placed on the developers, such as specific hardware to be used, databases, or coding standards [48], [49]. * **External Interfaces:** Rules defining how the software interacts with users (GUI), external hardware, and other software systems (e.g., generating XML export files) [50].

Part 2: IEEE 830 Standard for SRS Structure

The **IEEE 830 standard** is the universally accepted template for organizing an SRS document [51].

Section	Component	Description
1.0	Introduction	Contains Purpose, Scope, Definitions, Acronyms, Abbreviations, References, and an Overview of the document [52], [53]. The scope defines exactly what the system will and will not do [53].
2.0	Overall Description	Contains Product Perspective (business case and external interfaces), Product Functions (high-level summary of capabilities), User Characteristics

Section	Component	Description
		(technical proficiency of end-users), Constraints, and Assumptions/Dependencies [54]-[55].
3.0	Specific Requirements	The most critical section detailing External Interfaces, Functional Requirements, Performance Requirements, Logical Database Requirements, Design Constraints, and Object-Oriented Models (if applicable) [55]-[49].

NOTE

> Functional requirements are broken down hierarchically using numerical IDs (e.g., 3.1.1, 3.1.2) for exact traceability [56], [57]. For example, "Subject Registration" (3.1) might be split into "Register Course" (3.1.1) and "Drop Course" (3.1.2) [57].

Lecture 18: Representation of complex programming logic

Part 1: Pitfalls of Bad SRS (The Narrative Essay)

One of the worst ways to document requirements is through an unstructured narrative essay [58]. When functional requirements, non-functional requirements, and constraints are randomly mixed into massive text paragraphs, several critical issues arise [58], [59]:

- **Difficult to Change:** Finding and updating a specific feature requires hunting through scattered text, often leading to missed updates [59], [60].
- **Ambiguity & Imprecision:** Long sentences become confusing, and reading isolated sentences can provide incorrect meanings [60], [61].
- **Contradictions:** In lengthy text documents, authors often forget what was previously written, leading to conflicting rules [61].
- **Noise & Silence:** Narrative essays often include irrelevant data ("Noise") while completely omitting critical system interfaces or rules ("Silence") [62].

- **Over-Specification:** Essay formats often try to dictate *how* to solve a problem (e.g., dictating that data must be stored in descending order in a file), improperly restricting the architect's design flexibility [63], [64].
- **Forward References & Wishful Thinking:** Referencing concepts not yet defined in the document, or promising functionalities that are technically impossible to implement [65], [66].

Part 2: Mathematical & Graphical Representation of Complex Logic

When processing logic involves highly complex conditional checks, text-based descriptions fail because it is incredibly easy to miss variable combinations or overlapping conditions [67]. To ensure clarity for both developers and testers, analysts must use semi-formal techniques to represent complex logic [68].

The two most popular techniques are **Decision Trees** and **Decision Tables** [68].

1. Decision Trees

A graphical representation where processing logic is broken down into a branching structure [69]. * **Nodes (Junctions):** Represent the conditions or decisions to be evaluated [69]. * **Edges (Branches):** Represent the outcome of the condition (e.g., True/False, Valid/Invalid) [69]. * **Leaf Nodes:** Represent the final action to be taken based on the pathway [69]. * *Advantage:* Provides an immediate, easy-to-understand visual map of logic flow, making it simple to spot missed conditions [70].

2. Decision Tables

A tabular matrix approach that pairs condition states with corresponding actions [71]. Decision tables are highly compact and far more scalable than trees when dealing with multiple intersecting variables [72].

Structure of a Decision Table: * **Top Half (Conditions):** Lists all condition variables (V1, V2, etc.) and their possible states (True/False or specific values) [71], [73]. * **Bottom Half (Actions):** Lists all possible actions (A1, A2, etc.) that the system can execute [74], [73]. * **Columns (Rules):** Each column represents a distinct "Rule" — a specific combination of conditions and the resulting action [75], [76]. Every column easily translates into a direct Test Case for quality assurance [77].

Example Breakdown Matrix (Decision Table Format):

Conditions / Rules	Rule 1	Rule 2	Rule 3
Selection Type	New Member	Update Member	Cancel Member
Valid Selection?	Yes	Yes	Yes
Action 1	Prompt for details	Update expiry date	Delete record
Action 2	Create record	Show success message	Print refund check
Action 3	Print bill	-	-

(Note: If the selection is invalid, all condition columns are ignored, and a generic error message is printed [78], [79].)

 IMPORTANT

> **Why Represent Logic Formally?** A decision table forces the analyst to calculate all possible combinations of conditions (2^n combinations for boolean variables), ensuring zero missing edge-cases [75], [76]. This mathematical rigor guarantees that developers build exactly what is required and testers have an explicit matrix to validate against [77], [79].

Lecture 19: Design Fundamentals

Part 1: The Software Design Process

The software design phase acts as the bridge between requirements engineering and implementation. It transforms the **Software Requirement Specification (SRS)** document into a structured blueprint that programmers can directly translate into executable code.

- **Input:** SRS Document

- **Process:** Iterative refinement, evaluating design alternatives.
- **Output:** Software Design Document (SDD).

Part 2: Procedural vs. Object-Oriented Design

Design fundamentals vary heavily based on the underlying programming paradigm chosen for the project: * **Procedural Design:** Focuses on functions and top-down decomposition.

The design elements include module structures, call relationships, parameter interfaces, local data structures, and algorithms. Modules consist of functions and global data. *

Object-Oriented (OO) Design: Focuses on real-world entities. The design elements include class structures, inheritance, object invocations, and message passing.

NOTE

> Regardless of the paradigm (Procedural or OO), the design process is highly iterative. It requires continuous intellectual stimulation and evaluating multiple design alternatives before finalizing the architecture.

Part 3: High-Level vs. Low-Level Design

The design phase is divided into two sequential stages:

1. High-Level Design (Architectural Design):

- **Focus:** Identifies the overall architecture, the modules required, and their interactions.
- **Output:** Module structure, inter-module control relationships (call invocations), and interfaces (data exchanged).
- **Representations:** Structure Charts, Jackson Diagrams, Warnier-Orr Diagrams.

2. Low-Level Design (Detailed Design):

- **Focus:** Zooms into the internal workings of each module identified in the high-level design.
- **Output:** Specific data structures, algorithms, and logical flows for individual modules.
- **Representations:** Module specifications, pseudocode, flowcharts.

Part 4: Characteristics of a Good Software Design

Evaluating whether a design is "good" or "bad" is critical when choosing between design alternatives. A robust design must exhibit: 1. **Correctness:** It must accurately and completely implement all functional requirements specified in the SRS. 2.

Understandability: The design must be clean and intuitive. Code spends the majority of its lifecycle in the maintenance phase; if maintainers cannot understand the design, updates become costly and error-prone. 3. **Maintainability:** The architecture must be resilient to change, allowing modifications without triggering cascading system failures. 4. **Efficiency:** Optimal use of computational resources.

Lecture 20: Modular Design

Part 1: Principles of Modularity

Modularity is the foundational principle for achieving **understandability** and **maintainability** in software design. It is the practice of breaking down a large, complex system into smaller, manageable, and nearly independent components called **Modules**.

IMPORTANT

> ****Divide and Conquer Principle:**** Modularity directly implements this principle. A massive system is impossible to grasp at once. By dividing it into independent modules, a developer can comprehend, test, and debug one module at a time without needing to hold the entire system in memory.

Benefits of Modularization:

- * **Bug Localization:** In a well-layered, tree-like module structure, errors are confined to specific branches. A bug in a leaf node doesn't arbitrarily crash unrelated higher-level modules.
- * **Reusability:** Independent modules can be extracted and reused in entirely different projects.
- * **Reduced Complexity:** Simplifies the cognitive load on developers.

Part 2: Module Layering and Tree Structures

A good modular design naturally forms a hierarchical, tree-like structure. * **Abstraction:** Higher layers represent broader abstractions, while lower layers handle concrete implementations. * **Low Fan-Out / Appropriate Fan-In:** A module should not control an excessive number of sub-modules (Low Fan-out), but heavily utilized utility modules can be called by many higher-level modules (High Fan-in). * **Avoid Cycles:** Circular dependencies (where Module A calls B, B calls C, and C calls A) create spaghetti code, making bug tracking impossible.

Part 3: Functional Independence

For modularity to be effective, modules must exhibit **Functional Independence**. This means a module performs a precise, well-defined function with minimal interaction with other parts of the system.

Functional independence is measured using two critical metrics: 1. **Cohesion:** How closely related the internal elements of a single module are. 2. **Coupling:** How heavily connected one module is to other modules.

TIP

> **The Golden Rule of Software Design:** Strive for **High Cohesion** and **Low Coupling**.

Lecture 21: Classification of Cohesion

Part 1: Understanding Cohesion

Cohesion measures the single-mindedness of a module. A highly cohesive module does exactly *one* thing and does it well. If you cannot describe what a module does in a single, simple sentence, it likely suffers from poor cohesion.

Part 2: The Cohesion Spectrum

Cohesion is classified into seven distinct levels. Evaluating a design involves categorizing its modules into these levels to assess quality.

Level	Cohesion Type	Quality	Definition & Characteristics	Example
1	Coincidental	Worst	A random collection of unrelated functions grouped into a single module with no logical basis.	A module named <code>AAA</code> that handles <code>print-inventory</code> , <code>register-student</code> , and <code>issue-book</code> .
2	Logical	Bad	Elements are grouped because they perform the same <i>logical</i> category of operation, even though they handle entirely different data.	A single module handling <i>all</i> printing: <code>print_quality</code> , <code>print_certs</code> , <code>print_salary</code> .
3	Temporal	Poor	Elements are grouped purely because they must execute in the same timeframe.	An <code>Initialize</code> module that sets up variables, opens files, and boots hardware all at startup.
4	Procedural	Fair	Elements are grouped because they follow a specific sequence of execution in an algorithm, though they	A module that computes a value and then immediately triggers a totally unrelated hardware sensor.

Level	Cohesion Type	Quality	Definition & Characteristics	Example
			may not share data.	
5	Communicational	Good	Elements are grouped because they operate on the same core dataset.	A module processing a student record: <code>search_student ,</code> <code>print_student ,</code> <code>save_student .</code>
6	Sequential	Better	The output of one internal function serves as the direct input to the next internal function.	<code>Sort_Data -></code> <code>Search_Data -></code> <code>Display_Data .</code>
7	Functional	Best	Every element in the module contributes exclusively to the execution of a single, well-defined task.	<code>Compute_Square_Root</code> or <code>Calculate_RMS .</code>

Lecture 22: Classification of Coupling

Part 1: Understanding Coupling

Coupling measures the degree of interdependence between two or more modules. While zero coupling is impossible (modules must communicate to form a system), the goal is to

keep communication as simple and sparse as possible.

💡 IMPORTANT

> **Why Low Coupling?** > Highly coupled modules behave like tangled wires. If you change a data structure in Module A, Module B breaks. If you want to reuse Module C in a new project, you are forced to drag along Modules D, E, and F because they are tightly coupled. Low coupling isolates changes, making the system modular and maintainable.

Part 2: The Coupling Spectrum

Coupling is categorized into five standard levels based on how data and control are passed between modules.

Level	Coupling Type	Quality	Definition & Characteristics	Example
1	Content	Worst	One module directly modifies or relies on the internal workings, local data, or code of another module.	Module A directly alters a private local variable inside Module B.
2	Common	Bad	Multiple modules share access to the same global data space.	Modules A, B, and C all read and write to a global <code>User_Session</code> array.
3	Control	Poor	One module passes a control flag to another, explicitly dictating its internal execution flow.	Module A passes an <code>is_error</code> boolean to Module B, forcing B to execute its error-handling branch.
4	Stamp	Good	Modules pass entire, complex composite data structures, even if the receiving	Passing an entire <code>Employee_Record</code> struct to a module that

Level	Coupling Type	Quality	Definition & Characteristics	Example
			module only needs a small piece of it.	only needs the <code>employee_ID</code> integer.
5	Data	Best	Modules communicate exclusively by passing simple, primitive data elements (like integers, floats, chars) via parameters.	Module A passes <code>base</code> and <code>height</code> (integers) to a <code>Calculate_Area</code> module.

Part 3: Trade-offs and Structural Dynamics

Designing an optimal system requires balancing cohesion and coupling via structural design patterns.

- **Top-Down Decomposition:** This approach naturally fosters high cohesion and low coupling. You take a macro-function (e.g., `Create New Library Member`) and decompose it into discrete sub-functions (`Create Record`, `Assign Member ID`, `Print Receipt`).
- **Structure Charts (Fan-in / Fan-out):**
 - **Fan-out:** The number of sub-modules a module directly calls. *High fan-out* indicates a module is trying to do too much (low cohesion, poor abstraction).
 - **Fan-in:** The number of parent modules that call a specific module. *High fan-in* is excellent, as it indicates a highly cohesive, highly reusable module (e.g., a standardized error-logging function used throughout the system).
- **The "Why" Behind the Golden Rule:**
 - We strive for **High Cohesion** so that when business requirements change (e.g., how taxes are calculated), the developer only needs to update *one* specific, isolated module.
 - We strive for **Low Coupling** so that the update to the tax module does not trigger a cascading chain of compilation errors and runtime bugs across the reporting, UI, and database modules. Functional independence is the ultimate key to scalable software architecture.

Lecture 23: Introduction to structured analysis and structured design

Part 1: Structured Analysis and Structured Design (SA/SD) Methodology

The **Structured Analysis and Structured Design (SA/SD)** approach is a traditional methodology utilized by developers to transition from requirements to an executable system [1]. * **Structured Analysis (SA)**: Focuses on analyzing functional requirements and breaking them down via **Top-Down Decomposition** [2-4]. It acts as an extension of understanding customer requirements, transforming the SRS (Software Requirements Specification) document into detailed, manageable functional components [4]. * **Structured Design (SD)**: Takes the decomposed functional model (DFD) and transforms it into an architectural module structure [4, 5]. This procedural design strategy dictates how the system is modularized, ultimately enabling the writing of code in procedural languages like C [5].

IMPORTANT

> The core transition in SA/SD is shifting from a problem-oriented functional model (Structured Analysis) to a solution-oriented architectural model (Structured Design) [4, 5].

Part 2: The "Why" Behind Functional Decomposition and Data Flow Modeling

Functional decomposition is required because analyzing an entire complex system at once is practically impossible [3]. By identifying high-level functions from the SRS and progressively breaking them down into precise sub-functions, developers can methodically conquer complexity [3].

Data Flow Modeling is essential for this because it illustrates how data moves through the system, capturing which data is consumed as input, processed, and produced as output [6, 7]. It maps out a clear **Data Flow** (as opposed to Control Flow), acting as the blueprint for how information is transformed stage by stage [6].

Lecture 24: Basics of Data Flow Diagrams (DFD)

Part 1: Core DFD Symbols and Notation

Data Flow Diagrams are highly intuitive graphical models because they utilize a minimal set of standardized symbols (specifically 5 core symbols) to map out complex logic [8, 9].

Symbol Shape	Designation	Notation Rules & Description
Rectangle	External Entity (Terminator/User)	Represents an external system, user, or entity interacting with the software [10]. They supply input data or consume output data [11].
Circle / Bubble	Process (Transformation)	Represents a function or action. MUST be named using a verb (e.g., "Search Book", not "Book Search") [12].
Arrow	Data Flow	Represents the directional flow of data between entities, processes, and data stores [11]. Data names must be labeled on the arrow (except when connecting directly to a data store) [13].
Parallel Lines	Data Store	Represents data at rest (files, databases). Processes interact with them to read or write information [14]. External entities never directly access data stores [14].
Parallelogram	Output	Specifically used to represent output generation produced by the system (e.g., a printed report) [9, 13].

Part 2: DFD Hierarchy and Levels

DFDs utilize a hierarchical top-down model to maintain clarity [15].

- **Context Level DFD (Level 0 DFD):** Represents the entire software system as a single process bubble [16-18]. It defines the system's boundary by displaying all interactions (inputs/outputs) with **External Entities** [17, 18].
- **Level 1 DFD:** Decomposes the single Level 0 bubble into its core sub-functions (typically 3 to 5 processes) [19]. External entities are generally not repeated at this level; the focus is solely on internal data flows and processes [20, 21].
- **Level 2 DFD (and beyond):** Further breaks down each Level 1 process into its own detailed DFD [22]. A standardized numbering scheme ensures traceability: Level 1 processes are numbered 0.1, 0.2, 0.3. A Level 2 decomposition of process 0.2 would be numbered 0.2.1, 0.2.2, etc. [23, 24].

NOTE

> Data flow arrows can represent both ****synchronous**** dependencies (e.g., Process B cannot start until Process A generates a specific number) and ****asynchronous**** flows (e.g., data rests in a data store until needed) [25, 26].

Lecture 25: Developing DFD Model

Part 1: Managing Decomposition and DFD Balancing

The top-down decomposition of a DFD continues until the processes are simple enough that detailed pseudo-code can be easily written for them [20, 27, 28].

Balancing DFDs: A fundamental rule of DFD creation is **Balancing**. When a parent bubble is decomposed into a lower-level DFD, all the input and output data flows of the parent bubble **MUST** exactly match the net input and output data flows of its child diagram [29].

IMPORTANT

> If Process `A` takes data `X` and outputs data `Y`, the Level 2 DFD depicting Process `A` must collectively accept `X` as its only external input and emit `Y` as its only external output [29].

Part 2: The Data Dictionary (DD)

Whenever a DFD model is created, a **Data Dictionary** must be attached [30]. It serves to standardize the nomenclature of data across the project, preventing team members from using inconsistent terminology [30, 31]. It explains exactly what a complex data element consists of, breaking it down into primitive data elements [31].

Data Dictionary Grammar/Notation Rules: `*` `+` : Composition / AND (e.g., `A = B + C` means A consists of B and C) [29, 32]. `*` `[ ]` : Selection / OR (exclusive choice). `*` `{ name }5` : Exact repetition (5 instances of `name`) [32]. `*` `{ name }*` : Iteration (0 or multiple instances of `name`) [29]. `*` `=` : Equivalence [29].

Lecture 26: Examples of DFD Model development

Part 1: Step-by-Step Walkthrough (Trading-House System)

To conceptualize real-world DFD application, consider the Trading-House Automation System [33, 34]:

- Level 0 (Context Diagram):**
 - Bubble:** "Trading-House Software" [35].
 - External Entities:** Customer, Manager, Purchasing Department [35, 36]. (Accounts department logic is automated internally) [37, 38].
 - Flows:** Customer inputs orders and receives material issue slips [36, 38]. Manager requests and receives statistical queries [36, 38]. Purchasing issues indents [36, 39].
- Level 1 DFD:**
 - Processes Identified:** `Receive-Order`, `Process-Order`, `Handle-Indent-Request`, `Handle-Query` [39-41].
 - Data Stores:** Customer File, Item File, Inventory, Pending Order File, Sales Statistics [39, 40].
 - Flow Example:** `Receive-Order` checks customer credit [39]. Valid orders flow to `Process-Order` [40]. If items are available, it updates `Inventory` and `Sales`

Statistics and outputs a material issue slip [40]. If unavailable, it updates the Pending Order File [40].

Part 2: Common Errors Checklist in DFD Development

Reviewing DFDs requires actively searching for specific structural rule violations [42, 43].

Checklist for Preventing DFD Errors:

- * [] **Multiple Context Bubbles:** The Context Diagram (Level 0) must contain exactly **one** bubble [21].
- * [] **External Entities in Lower Levels:** External entities must only appear in the Context Diagram [21].
- * [] **Density Limits:** Any single DFD level should restrict itself to 3 to 7 bubbles to prevent cognitive overload [44].
- * [] **Control Flow vs. Data Flow:** Arrows must only represent data. Drawing a "control arrow" (an arrow without data just to show execution sequence, like pointing from "Validate Number" to "Generate Error") is strictly incorrect [42, 44-47].
- * [] **Direct Data Store Connections:** Data cannot flow magically on its own. Direct flows between a Data Store and an External Entity, or between two Data Stores (e.g., Inventory to Item-File), are invalid. A process bubble must always mediate [46, 48].
- * [] **Unbalanced DFDs:** Parent input/outputs must match child input/outputs [43].
- * [] **Missing Nomenclature:** Forgetting to explicitly name the data flowing on an arrow [43].
- * [] **Miracles and Black Holes:** (Conceptually derived) A bubble that consumes data but outputs nothing, or a bubble producing output without any data input. Every process must have at least one input and one output.
- * [] **Noun Process Names:** Bubbles named with nouns instead of verbs [49].
- * [] **SRS Mismatches:** DFD containing extra functions not in the SRS, or missing functions mandated by the SRS [43, 47].

Lecture 27: DFD Model - More Examples

Part 1: Limitations of the DFD (Analysis Model)

Data Flow Diagrams (DFDs) provide a simple, elegant method for identifying various processes at different levels and mapping the data exchanged between them [1].

However, relying solely on DFDs for system design introduces significant drawbacks: *

Ambiguity and Guesswork: A DFD bubble (process) only provides a high-level name.

Developers are left to guess the exact internal processing logic, leading to ambiguity and incomplete understanding [2]. * **Lack of Control Flow:** DFDs only show how data moves through the system; they do not explicitly define the sequence of execution or control logic (e.g., loops or conditional branches) [1, 3]. * **Requirement for Data Dictionary:** To mitigate ambiguity, a DFD must always be supported by a comprehensive Data Dictionary to clarify the exact definitions of data elements [3].

IMPORTANT

> **"The 'Why' Behind Converting DFDs to Structure Charts"** > DFDs are strictly an **"analysis model"** used to understand and validate user requirements [1, 3]. To move toward actual implementation, developers must transition to a **"design model"** called a **"Structure Chart"**. The Structure Chart transforms abstract data flows into a concrete module hierarchy with explicit control flows and invocation sequences, serving as the direct blueprint for writing code [3, 4].

Part 2: Essentials of the Structure Chart

The Structure Chart represents the high-level architectural design of the software [4]. It breaks down the system into a set of fundamental modules and establishes their hierarchical relationships.

Core Symbols in a Structure Chart:

1. **Module:** Represented by rectangular boxes [4].
2. **Control Call / Invocation:** Represented by a standard arrow pointing from a higher-level module to a lower-level module, indicating that the upper module "calls" the lower one [4, 5].
3. **Selection (Diamond):** A diamond symbol on a call arrow indicates a conditional call (the module is invoked only if a certain condition is met) [5].
4. **Repetition / Loop:** A curved arrow indicates that the module is called repeatedly (e.g., within a loop) [5].

TIP

> **"Design Heuristics:"** A well-designed Structure Chart strictly follows a top-down hierarchy. **"Backward arrows"** (arrows pointing from a lower-level module back to a higher-level module) are strictly prohibited. Such cyclical relationships violate abstraction principles, making the design overly complex, difficult to understand, and extremely hard to debug [5, 6].

Lecture 28: Essentials of Structure Chart

Part 1: Evaluating Structure Chart Quality

A good structured design must adhere to strict layering principles. A structure chart should resemble a clean tree hierarchy [7, 8]. * **Good Design:** Exhibits clear, distinct layered levels with clean top-down invocation [8]. * **Bad Design:** Skips layers or uses backward-pointing arrows. Such designs fail to maintain abstraction and create chaotic dependencies [8].

Part 2: Transform Analysis

Transform Analysis is the methodical process of converting a DFD into a layered Structure Chart [9]. To map the DFD successfully, the system is conceptually divided into three core branches:

1. **Input Branch:** This branch represents processes that read data from the user and validate or filter it [10, 11]. The data here is fundamentally driven by the input source [12].
2. **Output Branch:** This branch takes logical data and formats, structures, or writes it to an output medium (e.g., table, tree, or display) [9, 10].
3. **Transform Center (Central Processing):** The remaining central processes make up the Transform Center. This is where the core business logic occurs [9, 10]. The Transform Center relies on pre-existing data or complex logic to convert the input into the logical output; it is not a simple pass-through mechanism [12].

NOTE

> ****Validation Processes**:** Processes that merely check or validate input data belong to the ****Input Branch****, not the Transform Center. However, if a process sorts or deeply filters the data using business rules, it may qualify as part of the Transform Center [11].

Step-by-Step Transformation Workflow: 1. Identify the Input, Output, and Transform Center portions of the DFD [10]. 2. Create a top-level root module to serve as the main controller (e.g., `Compute-RMS`) [13]. 3. Draw first-level factored modules corresponding to the input (e.g., `Get-Good-Data`), transform, and output branches [13]. 4. Continue

factoring and mapping lower-level DFD bubbles into sub-modules in the Structure Chart until the design is complete [13, 14].

Lecture 29: Structure Chart Development

Part 1: Transform vs. Transaction Analysis

While Transform Analysis is ideal for linear, sequential processing, **Transaction Analysis** is required when the system handles multiple distinct types of operations based on a single input [15, 16].

- **Transform Analysis Example (Tic-Tac-Toe):** In a Tic-tac-toe game, the human player inputs a move. The system reads the move, validates it, and updates the board data structure. Because this is a straightforward sequence of events, it is easily mapped using Transform Analysis [17].
- **Transaction Analysis (Menu-Driven Systems):** In many applications, a single data item triggers multiple completely different processing paths (e.g., a software package menu). The DFD will show an input splitting into several different transactions [15, 16].

Part 2: Factoring and Real-World Application

When applying Transaction Analysis, the top-level Structure Chart will feature a **Transaction Center** module. This controller module evaluates the input and dispatches the execution to one of several distinct lower-level modules (the transaction paths) [16].

- **Library Management Example:** A Library Management software cannot be modeled with a simple linear DFD because it handles multiple distinct transactions (e.g., issue book, return book, register member). In such cases, developers must move beyond the Context Diagram to the Level 1 DFD to identify the different transaction branches and apply Transaction Analysis to factor them into separate modules [18].

Procedural design is inherently a **top-down process**. Developers start with the high-level requirements, decompose them into DFDs, and apply these design heuristics to factor out

complex processes into manageable, codable Structure Charts [19].

Lecture 30: Structured Design Examples

Part 1: Paradigm Shift to Object-Oriented Design (OOD)

While previous lectures focused heavily on procedural design and top-down decomposition, the industry standard has largely shifted toward **Object-Oriented Design (OOD)** [20].

IMPORTANT

> **The Standardization of UML** > In the early 1990s, numerous OOD methodologies (such as Rumbaugh's OMT) existed, each with its own proprietary notations. This lack of standardization made it nearly impossible for different organizations or student teams to understand and reuse each other's designs [21, 22]. In 1997, the **Unified Modeling Language (UML)** was introduced to standardize these notations into a single, cohesive framework, borrowing heavily from successful methods like OMT while introducing new standards [22, 23].

Part 2: Overview of UML Modeling

UML allows software architects to capture different perspectives of a system through a unified vocabulary and standardized diagrams [24]. The core perspectives include:

1. **User View (Behavioral)**: Represented by the **Use Case Diagram**. Because the user's perspective forms the absolute starting point for all subsequent design, it is placed at the center of the UML modeling process [24, 25].
2. **Structural View**: Represented primarily by the **Class Diagram**, which outlines the static architecture of the system [25].
3. **Interaction/Communication View**: Represented by Interaction Diagrams (such as Sequence Diagrams), which focus on the dynamic behavioral exchange of data and

messages between objects [25].

(Note: Real-world examples of Supermarket Automation and Trading Systems are detailed in other segments of the course methodology, specifically covering advanced Use Case Modeling and earlier DFD basics, respectively, but the foundational mapping heuristics established in these lectures apply universally across all such domains).

Lecture 31: Use Case Modelling

Part 1: Core Concepts of Use Case Modelling

The **Use Case Model** is the central diagram in Object-Oriented (OO) design and UML modelling because it captures the system from the user's perspective [1, 2]. Although part of OO development, a use case model is technically a functional model rather than a purely object-oriented one, as it captures the functionalities the system must perform rather than the objects themselves [3, 4].

IMPORTANT

> A **Use Case** corresponds to a high-level requirement [4]. It provides a graphical representation of the system's intended behavior, making it easier to understand who uses the system and what operations they perform [5, 6].

The "Why" behind this modeling is to provide a clear, graphical view of all requirements, which aids in designing the Graphical User Interface (GUI) and tailoring user manuals to specific user proficiency levels [7-9].

Part 2: Actors and System Boundaries

Actors are the users of the system [4]. They can be human users (e.g., Library Member, Librarian, Accountant) or **External Systems** that interact with the software (e.g., a backup system or a credit card authorization system) [4, 10, 11]. * **Primary Actors** initiate the use case (e.g., a customer placing a telephone order) [12, 13]. * **Secondary/External Actors** are invoked during the use case (e.g., an external system verifying credit) [11, 14].

Use Case Diagram Elements:

- * **Use Case:** Represented by an ellipse/eclipse shape labeled with a verb phrase (e.g., "Issue Book") [15, 16].
- * **Actor:** Represented by a stick figure icon [9].
- * **Association:** A solid line connecting the actor to the use case, indicating that the user invokes the functionality [9, 17].
- * **System Boundary:** A rectangle enclosing the use cases, indicating the scope of the system or a specific version of it [17]. It is optional but recommended for clarity [16].

Lecture 32: Factoring Use Cases

Part 1: Why Factor Use Cases?

Factoring Use Cases means breaking down complex use cases into simpler, more manageable ones [18].

NOTE

> There are two primary reasons to factor use cases:

1. **To Simplify Design:** Complex use cases lead to overly complicated sequence diagrams and design models [18, 19].
2. **To Prevent Duplication:** Factoring extracts common functionalities shared across multiple use cases, reducing design and coding effort [19, 20].

Part 2: Techniques for Factoring Use Cases

UML provides three main techniques for factoring use cases: **Generalization**, `<<include>>`, and `<<extend>>` [21].

1. Generalization-Specialization

Used when multiple use cases share a common base functionality but have slight variations [21]. The common logic is kept in the **Parent Use Case**, and variations are placed in the **Child Use Cases** [21, 22].

- * **Visual Representation:** A solid line with a hollow triangle pointing to the parent use case [23].
- * *Example:* "Pay Membership Fee"

(Parent) generalized into "Pay via Credit Card" and "Pay via Library Pay Card" (Children) [24, 25].

2. The `<<include>>` Relationship

Used when a base use case requires a common, mandatory functionality that is extracted into its own use case [26, 27]. * **Visual Representation:** A dotted arrow with the

`<<include>>` stereotype pointing *from* the base use case *to* the included use case [26]. *

Example: Both "Issue Book" and "Renew Book" must invoke "Check Reservation" [28, 29].

3. The `<<extend>>` Relationship

Used for optional functionalities or variations that extend the behavior of a base use case under specific conditions [29, 30]. * **Visual Representation:** A dotted arrow with the

`<<extend>>` stereotype pointing *from* the extending use case *towards* the base use case

[30, 31]. UML also allows defining an **Extension Point** in the base use case [32]. *

Example: "Perform Sale" can be optionally extended by "Gift Wrap Product" if the customer requests it [32].

Feature	<code><<include>></code>	<code><<extend>></code>
Execution	Mandatory (Base use case cannot complete without it) [30].	Optional (Triggered only under specific conditions) [30].
Direction of Arrow	Points towards the included use case [26].	Points towards the base use case [30].
Primary Goal	Reuse common functionality across multiple use cases [27].	Add optional/exceptional behavior to a base use case [30].

Lecture 33: Overview of Class diagram

Part 1: Documenting Use Cases

While graphical models show relationships, they do not provide the exact sequence of interactions [33]. UML does not strictly mandate a documentation format, but standard templates (like Alistair Cockburn's) are highly recommended [34].

**TIP**

> **Use Case Description Template Elements [34-37]:**

- > **Name:** Name of the use case.
- > **Actors:** The users or external systems involved.
- > **Trigger:** The event that starts the use case.
- > **Pre-conditions:** What must be true before the use case can execute (e.g., ATM must be in a 'ready' state and have cash).
- > **Post-conditions:** What must be true after successful execution (e.g., Account balance is updated).
- > **Main Flow (Mainline Sequence):** The standard, successful sequence of interactions between the actor and the system.
- > **Alternative Flows (Exceptions):** Variations or errors (e.g., wrong PIN entered, insufficient funds, power failure) [38, 39].

Part 2: Introduction to Class Diagrams

Class Diagrams are the backbone of OO modeling. They group entities with common features into **Classes** [40].

- * **Visual Representation:** A solid-outline rectangle divided into compartments: **Name**, **Attributes**, and **Operations** [40, 41].
- * **Progression in Design:** During early design stages, only the class name is identified. As design progresses, operations are added, and finally, attributes are defined for coding [42, 43].
- * **Naming Convention:** Class names should always be singular (e.g., "Book", not "Books") [41].

Lecture 34: Inheritance relationship

Part 1: The "Why" Behind Object-Oriented Relationships

In any OO software, classes do not exist in isolation; they relate to one another to build complex systems [44, 45]. Understanding the relationships between classes (Inheritance, Association, Aggregation/Composition, Dependency) is crucial for creating modular, reusable, and maintainable architectures [46].

Part 2: The Inheritance (IS-A) Relationship

Inheritance is a powerful OO mechanism that denotes a generalization-specialization (or **IS-A**) relationship [47, 48]. It allows a new derived class (child) to be defined based on an existing base class (parent) [47].

- **Code Reuse:** Derived classes automatically inherit the attributes and methods of the base class, preventing code duplication [49].
- **Overriding (Polymorphism Context):** Derived classes can redefine existing base class methods [50].
- **Abstracting Commonality:** Common attributes (e.g., Library Member ID) and operations (e.g., Issue, Return) are kept in the base class (Library Member), while specific additions are kept in derived classes (Student, Faculty, Staff) [51].

IMPORTANT

> ****Multiple Inheritance:**** A class derived from two or more base classes. For example, a "Research Student" class might inherit from both "Student" and "Staff" base classes [52].

Implementing Inheritance

In Java, inheritance is implemented directly using the `extends` keyword [53]. A constructor in a derived class (e.g., `Box`) will implicitly utilize the attributes (length, width) inherited from the base class (e.g., `Rectangle`) alongside its own new attributes (height) [54, 55].

Avoiding Incorrect Generalization

Inheritance should only be used when an **IS-A** relationship genuinely exists [56]. * *Wrong:* Deriving `Chassis`, `Engine`, and `Door` from a `Car` base class [57]. An engine *is not* a car; it is *part of* a car (this is Composition/Aggregation) [56]. * *Right:* `UG Member` **IS-A** `Library Member` [58].

Note on Pitfalls: Creating deep, complex inheritance hierarchies can lead to poor cohesion and strong coupling, making leaf classes incredibly difficult to understand and maintain [59].

Part 3: Comparing Relationships

Feature	Inheritance (Generalization)	Association
Concept	Represents an "IS-A" relationship (e.g., Dog is an Animal) [48, 49].	Represents a "Uses-A" or peer relationship (e.g., Person works for Company) [60, 61].
Implementation	Achieved via language keywords (e.g., <code>extends</code> in Java) [53].	Achieved by passing object references or storing objects as attributes.
Lifecycle	Static and compile-time bound.	Dynamic; links between objects can be established or dissolved at runtime [62].

Lecture 35: Association relationship

Part 1: Understanding Association

An **Association** is a structural relationship that specifies that objects of one class are connected to objects of another class [1]. It is typically represented as a solid line connecting two classes, sometimes utilizing an arrow to indicate navigability [2].

- **Links vs. Associations:** An association defines the static relationship between two classes, whereas a **link** is a dynamic, runtime instance of an association existing between specific objects [3]. Links can be created and dissolved dynamically during execution, whereas associations remain permanent in the static class structure [3].
- **Unary Association:** An association can also be defined within the same class, indicating that objects of a class interact with other objects of the same class (e.g., a

Person object being a friend to another *Person* object, or a *Computer* node connecting to other *Computer* nodes) [4, 5].

Part 2: Multiplicity and Role Names

Multiplicity defines how many instances of one class can be associated with a single instance of another class [2]. * **1 (Exactly one)**: E.g., A person has exactly one mother [6]. * **0..*** (**Zero or many**): Represented by an asterisk `*`. E.g., A woman can be the mother of many persons [6]. * **1..*** (**One to many**) / **Custom ranges**: E.g., A student can take **1..5** courses, and a course can have **10..300** enrolled students [7, 8].

Role names can be added to the association ends to clarify the nature of the relationship (e.g., a *Person* class and a *Company* class connected by an association named "works for," with roles like *Employee* and *Employer*) [9, 10]. Implementing this in code generally translates the roles into attribute types within the respective classes [11].

Part 3: Navigability (Bi-directional vs. Uni-directional)

- **Bi-directional Association**: Indicated by a simple solid line or a line with arrows on both ends [12, 13]. This means both classes are aware of each other and can traverse the relationship in both directions.
- **Uni-directional Association**: Indicated by an arrow at one end of the association line [12]. For example, if a *Key* opens a *Door*, the *Key* object can invoke a method on the *Door* object, but the *Door* does not invoke methods on the *Key* [12].

TIP

> Always provide a "reading direction" or an arrow when naming an association to prevent misinterpretation of the relationship's context [8, 13].

Lecture 36: Aggregation/ Composition and

dependency relations

Part 1: Aggregation vs. Composition

Both Aggregation and Composition are specialized forms of association representing a whole-part relationship, but they differ significantly in ownership and lifetime binding [14].

- **Aggregation (HAS-A / Weak Ownership):** Represented by an **open (empty) diamond** at the "whole" end [15, 16]. It denotes a relationship where the component parts can exist independently of the composite class [17]. For example, an organization employs members; if the organization ceases to exist, the members still exist independently [18].
- **Composition (PART-OF / Strong Ownership):** Represented by a **solid (filled) diamond** [15, 16]. It implies a strict lifetime dependency where the composite class assumes total responsibility for the creation and destruction of its parts [19, 20]. For example, a car and its wheels; when a car object is instantiated, 4 wheel objects must be created via its constructor, and destroying the car inherently destroys the wheels [20, 21]. An object can only be part of one composition at a time [20].

Why choose Aggregation vs. Composition vs. Association?

- **Use Association** when classes interact as peers without any clear "whole-part" hierarchy (e.g., a Doctor and a Hospital where the doctor can leave and join elsewhere) [22, 23].
- **Use Aggregation** when there is a whole-part structure, but parts are added/removed dynamically and outlive the whole (e.g., adding items to a restaurant order iteratively) [24].
- **Use Composition** when the parts are intrinsically bound to the whole and cannot exist without it (e.g., an order created strictly simultaneously with its specific immutable items, or a window creating its sub-frames) [20, 25, 26].

Feature	Aggregation	Composition
Symbol	Open/Empty Diamond	Solid/Filled Diamond
Relationship	Weak "HAS-A"	Strong "PART-OF"

Feature	Aggregation	Composition
Lifetime Dependency	Parts can exist independently of the whole	Parts are destroyed when the whole is destroyed
Sharing	Parts can be shared among multiple wholes	A part belongs to exactly one whole at a time

Part 2: Dependency Relation

A **Dependency** relationship signifies that one class (the dependent) relies on another class (the independent) [27]. * **Symbol:** Represented by a **dotted arrow** pointing from the dependent class to the independent class [27]. * **Meaning:** If a change is made to the independent (concrete) class, it will affect the dependent class [28]. It is commonly used when one class delegates responsibilities to another, uses another class as a parameter, or temporarily instantiates it [28, 29].

NOTE

> Often, stereotypes like `<>` or `<>` are attached to dependency arrows to clarify the exact nature of the reliance [28].

Feature	Association	Dependency
Nature	Structural (objects retain references to each other)	Behavioral / Usage-based (often temporary)
Symbol	Solid line	Dotted arrow
Lifetime	Static and continuous during object lifetimes	Transitory (e.g., passed as a parameter)

Lecture 37: Interaction Modelling

Part 1: Concepts of Interaction Modeling

Interaction modeling is crucial for understanding the **behavioral view** of a system, shifting focus from static class structures to dynamic runtime behaviors [30, 31]. * It models how objects collaborate and communicate to realize the functionality of a specific Use Case [31]. * A key rule of thumb is that the **number of interaction diagrams typically equals the number of use cases** identified in the system, as each diagram captures the execution flow of one use case [32, 33].

Part 2: Types of Interaction Diagrams

There are two primary types of interaction diagrams in UML that depict these message exchanges:

1. **Sequence Diagrams:** Focus heavily on the time-ordering of messages [34].
2. **Communication/Collaboration Diagrams:** Focus heavily on the structural organization and relationships of the objects interacting [34].

IMPORTANT

> Sequence diagrams are indispensable because they directly aid in method extraction. By observing the messages an object receives during an interaction, designers can deduce the exact methods that must be populated into the corresponding class [35-37].

Lecture 38: Development of Sequence diagrams

Part 1: Core Components of Sequence Diagrams

A Sequence Diagram is a two-dimensional chart where objects are arranged horizontally at the top, and time progresses vertically downwards [38, 39]. * **Anonymous Objects:** Represented as `:ClassName` with an underline (e.g., `:Book`) [40, 41]. * **Lifelines:** A vertical dotted line dropping from each object, representing its existence over time [39]. * **Activation Bars:** A small vertical rectangle drawn over a lifeline, denoting that the object is currently active and possesses execution control [42, 43]. * **Messages**

(Synchronous/Asynchronous): Represented by solid horizontal arrows between lifelines. The message name maps directly to a method call [43]. * **Return Messages:** Represented by **backward dotted arrows**. Often omitted to reduce clutter, but used explicitly when a specific return value drives subsequent logic [44-46]. * **Object Creation & Destruction:** Object creation is shown by a dotted arrow stereotyped with `<<create>>`, pointing to the object box placed lower on the timeline [47]. Destruction is represented by a large 'X' at the end of the lifeline [43, 45].

Part 2: Combined Fragments and Logic

UML incorporates interaction frames to represent complex programming logic directly within sequence diagrams [48]. * **Conditions (opt/alt):** Mutually exclusive paths or conditional messages are represented using square brackets `[condition]`. The message is only transmitted if the condition evaluates to true [44, 48, 49]. * **Loops (loop):** Iterations are traditionally denoted by an asterisk `*` placed before the message, indicating that the message is repeatedly sent [48, 50].

Part 3: Communication (Collaboration) Diagrams

A Communication diagram (formerly Collaboration diagram) is essentially another view of a Sequence Diagram and is automatically generated by most CASE tools [51, 52]. *

Structure: It omits lifelines and time dimensions, mapping objects and their links directly [53]. * **Sequence Numbers:** Because vertical time ordering is absent, messages are prefixed with hierarchical sequence numbers (e.g., 1, 1.1, 2, 3) to strictly define the order of execution [53, 54]. * **Purpose:** The primary goal of a communication diagram is to clearly highlight the **Class Associations**. If a message passes between two objects, an association line must structurally exist between them [52, 55].

Lecture 39: State-Machine diagram

Part 1: The "Why" Behind State Modeling

Interaction diagrams (Sequence and Collaboration diagrams) are crucial for identifying methods within classes, commonly known as **method population** [1]. While interaction diagrams document how objects interact to realize a use case, they do not easily capture the internal lifespan and state changes of complex objects. This is where **State-Machine Diagrams (Statecharts)** come into play. When a class exhibits significant state-dependent behavior, state modeling is required to properly map out transitions and automatically generate skeleton code [2].

Part 2: Limitations of Finite State Machines (FSM)

A traditional **Finite State Machine (FSM)** is a "flat" state machine consisting of simple states and transitions driven by events [3, 4]. However, traditional FSMs suffer from two major drawbacks in object-oriented software engineering: 1. **The State Explosion Problem**: The number of states grows exponentially as state variables increase [3, 5]. 2. **Lack of Concurrency**: FSMs only allow sequential execution; they cannot model systems where multiple states are active simultaneously [3].

IMPORTANT

> **State Explosion Example (Robot Modeling):** > Consider modeling a robot. It has a movement switch (on/off = 2 states), movement direction (forward, backward, left, right = 4 states), left arm (up/down = 2 states), right arm (up/down = 2 states), head direction (straight, left, right = 3 states), etc. [6-9]. > In a flat FSM, representing this requires $2 \times 4 \times 2 \times 2 \times 3 \times 2 \times 2 = 384$ distinct states [9]. This becomes impossible to comprehend or maintain.

Part 3: UML Statechart Diagrams

UML State Machine Diagrams solve the state explosion problem through **abstraction** and **decomposition** mechanisms [6]:

- * **Hierarchical States**: A state can contain nested sub-states, creating a hierarchy that simplifies the high-level view [4].
- * **Concurrent States**: Modeling independent state variables simultaneously (e.g., the robot's arm moving *while* its wheels move forward) drastically reduces the total number of explicitly drawn states [10, 11].

Part 4: Code Generation from State Models

State machines can be directly translated into executing code. A common approach is the **doubly nested switch in a loop** technique [12, 13]. * The outer `switch` block evaluates the **current state**. * The inner `switch` block evaluates the **event** occurring in that state. * The execution logic performs the transition and sets the new state [13, 14].

State Transition Table Example

Current State	Event (Trigger)	Guard Condition	Action (Do/Entry/Exit)	Next State
s1 (Idle)	e1 (Turn On)	[power_available == true]	do: initialize()	s2 (Active)
s2 (Active)	e2 (Move Forward)	[obstacle == false]	entry: startMotors()	s3 (Moving)
s3 (Moving)	e3 (Stop)	None	exit: stopMotors()	s2 (Active)

Lecture 40: An Object-Oriented design process

Part 1: The Object-Oriented Design Process Flow

The Object-Oriented Design Process takes the conceptual requirements and refines them into a specific implementation architecture [15]. The sequential flow (aligned with Unified Process principles) is as follows: 1. **Software Requirements Specification (SRS)** [16] 2. **Use Case Model**: Defines the user's view and high-level requirements [17]. 3. **Domain Model (Conceptual Class Diagram)**: Extracts the initial conceptual classes from the use cases [17]. 4. **Interaction Diagrams**: Maps the dynamic behavior between domain objects [17]. 5. **Class Diagrams**: Finalizes classes, attributes, and methods [17].

 NOTE

> While Use Case models are part of the OO design process, they are strictly functional models (capturing system behavior) and not true object-oriented models themselves [18, 19].

Part 2: Domain Analysis & MVC Architecture

Domain modeling involves identifying three primary types of objects from the SRS and Use Cases, heavily mirroring the **Model-View-Controller (MVC)** pattern [20]:

1. **Boundary Objects (View):** These represent user interfaces (GUI) [20, 21]. They gather input from the user and display output [21]. They contain **no business logic** [22].
2. **Entity Objects (Model):** These represent long-term persistent information [22]. Examples include a `Book` object (storing Name, Author, ISBN) or a `Book Register` [22, 23]. They act as "dumb servers" whose sole purpose is to store, retrieve, and view data [24].
3. **Controller Objects (Controller):** These are the overall coordinators responsible for realizing a use case [25]. A controller receives a request from a boundary object, executes the **business logic**, and delegates tasks to the appropriate entity objects [25-28].

Part 3: Identifying Domain Objects

- **Boundary Objects:** Derived directly from the Use Case Diagram. A boundary object is required for every **Actor-Use Case pair** [29, 30]. (e.g., 2 actors interacting with 1 use case = 2 boundary objects).
- **Controller Objects:** Generally, one controller is assigned per Use Case to manage its logic [31]. Complex controllers can be broken down further.
- **Entity Objects:** Identified through **Noun Analysis** of the problem description [32].

 TIP

> ****Noun Analysis Filtering:**** When identifying Entity objects, highlight all nouns in the SRS. Eliminate: > * Actors/Users (unless the system needs to maintain historical data on them, making them a "Surrogate User Class" like `Library Member`) [33-35]. >

* Passive verbs masking as nouns (e.g., "training", "approval") [34, 36]. > * Attributes (e.g., "Customer ID" is an attribute of a `Customer`, not an entity itself) [37].

Lecture 41: Domain Analysis

Part 1: Refining Object Extraction

Domain Analysis bridges the gap between the functional Use Case model and the structural Class model [38].

Case Study: Supermarket Prize Scheme * **Context:** 4 actors interacting with 3 use cases [39]. * **Boundary Classes (4):** `Register-Customer-Boundary`, `Clerk-Register-Boundary`, `Sales-Clerk-Sales-Boundary`, `Manager-Winner-Boundary` [39, 40]. *

Controller Classes (3): `Register_Customer_Controller`, `Register_Sales_Controller`, `Select_Winner_Controller` [40, 41]. * Controllers coordinate the actions of entity classes to deliver the final use case outcome [41].

Part 2: Advanced Noun Analysis Rules

When parsing the problem description for Entity objects, remember that Entity objects typically manage bulk records and instantiation [42, 43]. * For instance, a `Sales History` entity object acts as a collection that manages individual `Sales Record` entities [42]. * Entity class names should be **singular**. Even if representing a collection of 10,000 students, the class is named `Student`, not `Students` [35].

Case Study: Tic-Tac-Toe Game * **Context:** Human player and computer take turns marking a 3x3 square [44, 45]. * **Noun Extraction:** "3x3 square" translates to a `Board` entity class [45]. * **Attributes:** The `Board` consists of squares, and each square has an attribute of being "marked" or "unmarked" [45, 46].

Lecture 42: Examples of object-oriented design

Part 1: Sequence Diagrams and Behavioral Mapping

With Domain Objects identified, the next step is creating **Sequence Diagrams** to map out the business logic and object interactions [47].

Tic-Tac-Toe Sequence Flow: 1. User initiates a move -> `Game-Move Boundary` detects it [48]. 2. Boundary passes the move to the `Game-Move Controller` [49]. 3. The Controller holds the business logic. It checks if the move is valid or ambiguous [49]. 4. If valid, the Controller registers the move on the `Board` entity [49]. 5. The Controller calls a `check-winner` method on the `Board` [50]. 6. If no winner, the Controller prompts the Computer to make a move [50].

Part 2: Class-Responsibility-Collaboration (CRC) Cards

As systems scale (e.g., dozens of classes and heavy data exchange), sequence diagrams become incredibly complex and difficult to track [51]. To manage this, developers use **CRC Cards** (Class, Responsibility, Collaboration) [52].

CRC Card Structure & Workflow

CRC cards help developers easily identify which class holds which responsibility and which objects it must communicate with before attempting to draw a sequence diagram [52].

- **Process:** Team members distribute CRC cards among themselves. They read the use case description, role-play the scenario, and fill out what their specific class must do (Responsibility) and who it must contact (Collaborator) [53, 54].

Class Name:	<code>Game-Move Controller</code>
Responsibility (What I do)	Collaboration (Who I need help from)

Class Name:	Game-Move Controller
Validate player's move	<code>Board</code> (to check if square is marked)
Register valid move	<code>Board</code> (to update state)
Check for winner / draw	<code>Board</code>
Prompt computer move	<code>Computer AI</code> / <code>Game-Move Boundary</code>

Part 3: Finalizing the Object-Oriented Design

Once sequence diagrams (backed by CRC cards) are complete, CASE tools can automatically generate the preliminary **Class Diagram** by extracting the messages sent to each class and assigning them as methods [55, 56].

IMPORTANT

> **Controller Optimization:** > During design refinement, if a Controller object is found to be doing trivial work (e.g., merely passing messages between the Boundary and Entity without adding complex logic), the Controller can be **eliminated** to simplify the architecture [56, 57]. Conversely, if a Controller is overwhelmed with excessive business logic, it should be split into multiple smaller controllers [56].

Lecture 43: Basic concepts in Testing-I

Part 1: Core Testing Terminology & Economics

Understanding the foundational terminology is critical for software testing. The terms **Error**, **Fault**, and **Failure** are often used interchangeably, but they have distinct meanings in software engineering [1].

- **Error:** A human action or mistake made by a programmer during software development [1].
- **Fault (Defect/Bug):** The manifestation of that error in the code or software [1].

- **Failure:** The observable incorrect behavior of the software when a fault is executed [1].

IMPORTANT

> **Testing vs. Debugging** > **Testing** is the process of executing a program with the intent of finding failures by comparing the actual output with the expected output [2]. **Debugging**, on the other hand, is the process of locating the exact fault (bug) that caused the failure and fixing it [3-5].

The Economics of Early Bug Discovery

The cost of fixing bugs increases exponentially as the software progresses through the development lifecycle [6]. Identifying and fixing bugs early (e.g., during the requirements or design phase) significantly reduces development costs [4, 6]. If a bug is caught immediately after it is introduced, it is much cheaper to fix than waiting for the testing or maintenance phase [4, 7].

Part 2: Verification and Validation (V&V)

Verification is the process of evaluating software to determine whether it satisfies the conditions imposed at the beginning of that phase, essentially checking if it matches the documented design or requirements [8]. Techniques include reviews, simulations, and unit testing [4].

Testing Levels Hierarchy

1. **Unit Testing:** The initial level of testing performed by the developers to check individual modules or components [4]. If unit testing is skipped, tracking down the exact location of a bug later becomes incredibly difficult [5].
2. **Integration Testing:** Focuses on combining unit-tested modules and testing them together to identify interface issues and verify that the combined units function correctly [5, 9].

NOTE

> While the prompt specifically requested details on Integration Testing strategies (Top-Down, Bottom-Up, Sandwich, Big-Bang) and Regression Testing, the provided

sources for these lectures only cover the fundamental purpose of Integration Testing without detailing its specific derivative strategies or Regression Testing.

Lecture 44: Basic concepts in Testing-II

Part 1: System and Acceptance Testing

Moving up the testing hierarchy, testing expands to evaluate the system as a whole.

1. **System Testing:** Determines if the fully integrated system meets its specified functional and non-functional requirements (SRS) [10, 11]. It includes specialized non-functional testing such as performance testing, response time evaluation, usability, stress, recovery, and infrastructure testing [11, 12].
2. **Acceptance Testing:** This is typically performed by beta customers or end-users. The software is handed over to the customer, and if it passes their criteria, they accept it; otherwise, it is rejected [10, 11].

Part 2: The Pesticide Effect and Testing Strategies

A crucial concept in testing is the **Pesticide Effect** [13].

TIP

> **The Pesticide Effect Explained** > Just as insects build resistance to a specific pesticide over time, software bugs "hide" from a repeated testing strategy [13, 14]. If you apply the same testing technique repeatedly, the bugs that escaped it initially will continue to escape [14, 15].

To overcome the Pesticide Effect, organizations must use multiple, diverse testing strategies (acting as different "bug filters") [15, 16]. Research shows that individual techniques (like reviews or specific tests) might only uncover about 30% of the defects [16, 17]. By chaining different techniques, the total number of latent defects is drastically reduced [17-19]. However, testing is a heuristic process; it reduces the number of bugs but cannot guarantee 100% bug removal [20].

Negative Test Cases

A well-designed test suite must include **Negative Test Cases**. These tests ensure that the system behaves gracefully and does not crash when provided with invalid inputs (e.g., entering letters instead of numbers) [21, 22].

Lecture 45: Basic concepts in Testing-III

Part 1: Test Suite & Test Case Design Principles

A **Test Suite** is a carefully designed collection of test cases used to evaluate a system [23].

A standard **Test Case** must explicitly define: * **Preconditions**: The required state of the system before the test [23, 24]. * **Test Input**: The specific data provided to the system [24]. * **Expected Output**: The anticipated behavior or result [24]. * **Postconditions**: The state the system should be left in after execution [24].

Test reports track the execution date, tester name, pass/fail status, and detailed analysis if a failure occurred (e.g., system crash, wrong output, infinite loop) [24, 25].

The Impracticality of Exhaustive Testing

Exhaustive testing (testing all possible input combinations) is practically impossible and would take years [26, 27]. Similarly, random testing is highly inefficient, as providing random data might repeatedly execute the same code paths while entirely missing others, leaving bugs undetected [27, 28]. Therefore, an optimal test suite must be intelligently designed to be reasonably sized while uncovering the maximum number of faults [27].

Part 2: Testing Methodologies & Environments

Testing strategies are broadly categorized into two visual analogies:

Feature	Black-Box Testing	White-Box Testing
Focus	Customer usage, inputs, and functional requirements [29].	Internal code structure, logic, and implementations [29].
Knowledge Required	No knowledge of internal code is required.	Deep knowledge of code and internal design [30].
Execution Order	Generally performed first to identify problematic areas [29].	Often guided by black-box results to thoroughly test problem areas [29, 31].

IMPORTANT

> **Drivers and Stubs in Unit Testing** > During unit testing, isolated functions cannot run on their own. Developers must write **Drivers** (code to call the function being tested) and **Stubs** (dummy code to simulate external functions called by the unit) [32].

Lecture 46: Unit testing strategies-I

Part 1: Grey-Box Testing and Black-Box Challenges

Bridging the gap between Black-Box and White-Box testing is **Grey-Box Testing** [33]. This approach is used particularly for components and classes, sitting intermediately between having zero internal knowledge and full internal knowledge [33, 34].

The Complexity of Black-Box Testing

Black-box testing relies on input data to observe output behavior [34]. However, testing multiple parameters causes a combinatorial explosion. For example, testing a simple

function checking equality between two 64-bit integers would require 2^{128} combinations, making exhaustive black-box testing impossible [35, 36]. Testing strategies must drastically reduce this number while remaining effective [37].

Part 2: Black-Box Strategies: Scenarios and Equivalence Classes

To optimize testing, two primary Black-Box strategies are used:

1. Scenario-Based Testing Based on Use Case models, testers identify the **Main Scenarios** and **Alternate Scenarios** [38]. Test cases are designed to trigger these specific paths through the software [39]. Each scenario demands specific pre-conditions, input values, and expected results [40].

2. Equivalence Class Partitioning This technique divides the vast input domain into smaller **Equivalence Classes** [41, 42]. The core principle is that the program will exhibit the same behavior for any input within a specific equivalence class [42].

- **Valid Equivalence Classes** (ஏற்றுக்கொள்ளக்கூடிய): Sets of acceptable, normal inputs [43, 44].
- **Invalid Equivalence Classes** (தவறான): Sets of inputs that fall outside acceptable bounds (e.g., negative numbers, letters instead of numbers) [43-45].

By selecting just one representative value from each valid and invalid equivalence class, testers can drastically reduce the number of test cases while ensuring broad coverage [42, 46].

TIP

```
> **Defining Equivalence Classes** > If a system accepts a range from 1 to 5000: > *  
*Valid Class:* Values between 1 and 5000. > * *Invalid Class 1:* Values < 1. > *  
*Invalid Class 2:* Values > 5000 [47, 48].
```

Lecture 47: Unit testing strategies-II

Part 1: Black-Box Testing & Equivalence Class Partitioning (ECP)

Black-box unit testing focuses on testing the functionality of a module based on its input-output behavior, without requiring knowledge of the internal code structure [1, 2]. Exhaustive black-box testing—testing all possible combinations of input values—is mathematically and practically impossible [3, 4]. For example, a simple function comparing two 64-bit integers would require $2^{64} \times 2^{64} = 2^{128}$ test cases, an astronomically large number [4, 5].

To overcome this, **Equivalence Class Partitioning (ECP)** is used. This strategy divides the input data domain into sets of data called **Equivalence Classes** [6].

IMPORTANT

> The fundamental assumption of ECP is that all input values within a single equivalence class will exhibit identical behavior and execute the same control flow path in the program [7, 8].

Valid vs. Invalid Equivalence Classes

For any given input, we must identify: * **Valid Equivalence Classes:** The set of inputs that are acceptable and represent normal operating parameters [9, 10]. * **Invalid Equivalence Classes:** The set of inputs that are unacceptable, out of bounds, or conceptually wrong [9, 10].

Example 1: Numeric Range If a function accepts an integer between 1 and 5000 [11]. * *Valid Class:* Integer values from 1 to 5000 [11]. * *Invalid Classes:* Values < 1 , and values > 5000 [12].

Example 2: Multiple Valid Classes If a function `issue-book` takes a Book ID [13]. * *Valid Classes:* Issuable book, Reference book (cannot be issued), Single-part book, Multi-part book [14, 15]. * *Invalid Classes:* Invalid Book ID format [13].

Part 2: Equivalence Testing Coverage Formulas

When a function takes multiple parameters, test cases must be designed by combining the equivalence classes of each parameter [16, 17].

1. Weak Normal Equivalence Class Testing Assumes that errors are rarely combinatorial. We simply need to ensure that every valid equivalence class of every parameter is covered by at least one test case [18]. * **Formula:** If Parameter 1 has m valid classes and Parameter 2 has n valid classes (where $m > n$), the number of test cases required is m [19, 20].

2. Strong Normal Equivalence Class Testing Tests every possible combination of valid equivalence classes across all parameters [21, 22]. * **Formula:** If Parameter 1 has m valid classes and Parameter 2 has n valid classes, the number of test cases required is $m \times n$ [19].

3. Strong Robust Equivalence Class Testing Expands upon strong normal testing by aggressively combining both valid and invalid classes [23]. * **Formula:** If Parameter 1 has n_1 valid and n_2 invalid classes, and Parameter 2 has m_1 valid and m_2 invalid classes, the total number of test cases required is $(n_1 + n_2) \times (m_1 + m_2)$ [23].

 NOTE

> Strong robust testing yields a significantly higher number of test cases and exposes the most errors, but it may become impractical as the number of parameters increases [24].

Lecture 48: Equivalence Class Testing-I

Part 1: Rules for Identifying Equivalence Classes

Identifying equivalence classes correctly requires practice and an understanding of the problem domain [25, 26]. Several heuristic guidelines help derive these classes based on input types:

Input Condition Type	Valid Equivalence Classes	Invalid Equivalence Classes	Example
Continuous Range	1 Valid Class (within the range)	2 Invalid Classes (below min, above max)	Range 1 to 100 . Valid: 50 . Invalid: <1 and >100 [27].
Set of Values (Menu/List)	1 Valid Class (belongs to the set)	1 Invalid Class (does not belong to the set)	Options {A, B, C} . Valid: A . Invalid: D [28].
Boolean	1 Valid Class (True/False)	1 Invalid Class (Non-boolean chars/strings)	Valid: True . Invalid: Control chars [29].
Multi-Part / Formatted String	1 Valid Class (Matches exact format)	Multiple Invalid Classes (Breaks various format rules)	6-char Password. Valid: 6 chars . Invalid: <6 , >6 , special chars [30, 31].

Part 2: Step-by-Step ECP Execution

- Analyze the Input Data Domain:** Determine if the input is a range, a set, or a formatted string [26, 27].
- Partition into Broad Classes:** Start by dividing the domain into one broad valid set and one or more invalid sets [32].
- Refine Classes:** Break down the valid and invalid sets into more granular equivalence classes based on specific behavioral triggers (e.g., a URL input can be partitioned by protocol: HTTP, HTTPS, FTP, or local file) [33-35].
- Select Representatives:** Pick a single representative data point from each identified valid and invalid class [36].

Lecture 49: Equivalence Class Testing-II

Part 1: Handling Multi-Parameter Complexities

Functions commonly take multiple parameters, and the behavior of the function depends heavily on how these parameters interact [37, 38].

TIP

> **Weak Testing vs. Strong Testing:** Weak testing is generally acceptable when assuming a "single fault" model (i.e., errors are caused by one parameter failing independently). Strong testing is necessary when parameters interact with each other to cause a failure [39, 40].

Example Analysis: Bank Interest Calculation Consider a function `Calculate-Interest(Deposit Amount, Days)`. * **Parameter 1 (Deposit Amount):** * Valid: \$ < 1,00,000\$, \$ \ge 1,00,000\$ [41]. * **Parameter 2 (Days):** * Valid: \$ \le 15\$ days, \$ 15 - 180\$ days, \$ 180 - 365\$ days, \$ 1 - 3\$ years, \$ > 3\$ years [42, 43]. * For **Strong Equivalence Class Testing**, we combine every class of Parameter 1 with every class of Parameter 2, resulting in a comprehensive test suite that catches cross-parameter dependencies [41, 44].

Part 2: Introduction to Special Value and Boundary Value Analysis (BVA)

Special Value Testing targets values where testers intuitively suspect a program might fail [43]. * **Common Hazards (Boundaries):** The most common locations for programmer errors (e.g., writing `<` instead of `<=`) are at the boundaries separating equivalence classes [45, 46]. * **Special Hazards:** Specific domain-related values (e.g., testing the year `2000` or `2004` specifically to check for Leap Year logic errors) [47].

When an input domain is a continuous range, boundaries exist between the valid and invalid equivalence classes [46].

Lecture 50: Special Value Testing

Part 1: Boundary Value Analysis (BVA)

Implementation

Boundary Value Analysis (BVA) mandates that test cases be explicitly designed around the edges of equivalence classes [48].

Why BVA? The Mathematical "Why"

Programmers frequently make logical errors at the extreme edges of input ranges—often due to off-by-one errors in loop counters or relational operator mistakes (using `<` when `<=` is required) [46]. BVA mathematically targets these precise failure points by isolating the boundary conditions.

BVA Test Case Selection

For any given input range $[A, B]$, boundary value testing requires selecting values at and immediately adjacent to the boundaries [49].

- * **Nominal (nom)**: A standard value comfortably inside the valid class [50].
- * **Minimum (min)**: The exact lower boundary [51].
- * **Minimum + 1 (min+)**: Just above the lower boundary [52].
- * **Maximum (max)**: The exact upper boundary [51].
- * **Maximum - 1 (max-)**: Just below the upper boundary [52].
- * *For Robustness*: **min-1** (just below the minimum, entering the invalid class) and **max+1** (just above the maximum, entering the invalid class) [50].

IMPORTANT

> **BVA Test Case Formulas:**

- > **Normal BVA ($4n + 1$ formula):** For n variables, we hold all variables at their nominal values while varying one variable at a time through its 4 boundary conditions (min, min+, max-, max). This gives $4n$ cases, plus 1 base case where all variables are nominal.
- > **Robust BVA ($6n + 1$ formula):** Expanding on normal BVA, we also test the invalid boundaries (min-1 and max+1). This adds 2 more cases per variable, yielding $6n$ cases, plus the 1 all-nominal case.

BVA Matrix Example (2 Parameters)

For 2 parameters (e.g., Age and Years of Education), applying the $4n+1$ logic results in $4(2) + 1 = 9$ test cases [53].

Test Case #	Parameter 1 (e.g., Age)	Parameter 2 (e.g., Education)	Focus
1	Nominal	Nominal	Baseline / Common element [53]
2	min	Nominal	P1 Lower Boundary
3	min+	Nominal	P1 Just Above Lower
4	max-	Nominal	P1 Just Below Upper
5	max	Nominal	P1 Upper Boundary
6	Nominal	min	P2 Lower Boundary
7	Nominal	min+	P2 Just Above Lower
8	Nominal	max-	P2 Just Below Upper
9	Nominal	max	P2 Upper Boundary

Part 2: Transition to Combinatorial Testing

As the number of parameters increases, combining all equivalence classes or boundary values (Strong Testing) leads to an unmanageable explosion of test cases (e.g., testing all interactions of 5 parameters) [54, 55]. To handle this, **Combinatorial Testing** strategies are introduced [55]. * **Decision Table-Based Testing:** Excellent for capturing complex business logic where specific combinations of inputs trigger specific actions [55, 56]. * **Cause-Effect Graphing:** Visually maps conditions (causes) to actions (effects) to systematically derive test cases [55]. * **Pairwise Testing:** A technique that drastically reduces test cases by testing all possible pairs of parameter values, assuming that most combinatorial errors arise from the interaction of just two variables [55].

Lecture 51: Combinatorial Testing

Part 1: The Combinatorial Explosion Problem

In software testing, verifying the behavior of an application across various input conditions is a critical task. However, as the number of input conditions increases, the number of possible test combinations grows exponentially [1]. This phenomenon is known as the **Combinatorial Explosion Problem**.

- For example, if a software component has 4 binary input conditions, testing all combinations requires $2^4 = 16$ test cases [1].
- If there are 10 binary conditions, it escalates to $2^{10} = 1024$ test cases [2].
- For 40 inputs, each with 3 possible values, the test suite blows up to an unmanageable 1.2×10^{19} combinations, making exhaustive testing impossible within a human lifetime [3].

IMPORTANT

> **Why Combinatorial Testing?** > To test complex logic comprehensively without running millions of test cases, combinatorial testing techniques (like Decision Tables and Pairwise Testing) are utilized to drastically reduce the test suite size while retaining a high fault-detection rate [4, 5].

Part 2: Introduction to Logic-Based Testing

When the system's output depends on complex logical relationships between multiple inputs, combinatorial methods are required. The inputs are systematically mapped to outputs through a formalized structure [6].

Lecture 52: Decision Table Testing

Part 1: Structure of a Decision Table

A **Decision Table** is a structured, tabular representation of complex programming logic that maps input conditions to corresponding actions [7, 8].

Core Components:

- * **Condition Stub:** The top rows list all possible input conditions or variables [8, 9].
- * **Action Stub:** The bottom rows outline the actions or outputs to be executed [8, 9].
- * **Rules (Entries):** Each column in the table represents a specific combination of condition values (a **Rule**) and corresponds to a specific **Test Case** [8].

Part 2: Rule Reduction and Simplification

When generating decision tables, creating a column for every possible combination (e.g., $2^3 = 8$ rules for 3 boolean conditions) often leads to redundancy [10].

TIP

> **Rule Simplification Algorithm:** If two rules result in the exact same action and only differ by the value of a single condition, that condition can be marked as a **"Don't Care"** condition [11, 12].

By merging columns with "Don't Care" states, the number of required test cases is significantly reduced [12, 13].

Example: Airline Meal Service

- * **Conditions:** Flight > 50% full? (C1), Ticket > \$3000? (C2), Domestic Flight? (C3) [14].
- * **Actions:** Serve Meal? (A1), Is it Free? (A2) [15].

Reduction: If the flight is domestic (C3 = Yes), meals are never free regardless of whether the ticket is > \$3000 or the flight is > 50% full. Therefore, C1 and C2 become "Don't Care" conditions for the Domestic rule, collapsing multiple rules into one and reducing the test suite size to just 4 optimal test cases [12, 13].

Lecture 53: Cause Effect Graphing

Part 1: Moving from Decision Tables to Graphs

While decision tables are highly effective, a system with 10 conditions requires $2^{10} = 1000$ initial columns, which is cumbersome to manually build and simplify [1]. **Cause-Effect Graphing** visually models the relationships between inputs (**Causes**) and outputs (**Effects**) before converting them into a decision table [16].

Part 2: Elements of Cause-Effect Graphs

- **Causes (Inputs):** The specific input conditions (e.g., Deposit Period, Deposit Amount) [16].
- **Effects (Outputs):** The actions or outcomes triggered (e.g., Interest Rate applied: 6%, 7%, 8%, or 9%) [17].
- **Boolean Logic:** Causes are linked to effects using Boolean logic (e.g., **AND** logic is used to connect a condition like "Deposit < 1 Lakh AND Term < 1 Year") [2].

NOTE

> *The following specific constraints and logic gates (NOT, Identity, Inclusive, Exclusive, One and Only One, Requires, Masking) are standard industry concepts in Cause-Effect Graphing, though they are not explicitly detailed in the provided source transcripts.*

- > * **Identity:** If cause occurs, effect occurs.
- > * **Negation (NOT):** If cause does NOT occur, effect occurs.
- > * **OR / AND:** Standard boolean logical combinators.
- > * **Mutually Exclusive (E):** At most one cause can be true.
- > * **Inclusive (I):** At least one cause must be true.
- > * **One and Only One (O):** Exactly one cause must be true.
- > * **Requires (R):** If one cause is true, another specific cause must also be true.
- > * **Masking (M):** One effect masks another effect from occurring.

Example: E-commerce Discount Logic

- * **Causes:** Purchase > \$2000 (C1), Used SBI Card (C2) [18, 19].
- * **Effects:** 10% Discount (A1), 15% Discount (A2), Additional 5% Discount (A3) [19].
- * **Graphing Benefit:** The graph identifies dependencies and mutually exclusive states (e.g., you cannot simultaneously have "Purchase < \$2000" and "Purchase > \$2000 excluding discounts"). These logically impossible combinations are excluded from the test cases instantly [20, 21].

Lecture 54: Pairwise Testing

Part 1: The "Why" Behind Pairwise Testing

In systems with massive input configurations (e.g., testing an Android app across different keyboards, screen layouts, and input methods), exhaustive testing reaches astronomical numbers (e.g., 172,800 combinations) [22, 23].

Research shows that software failures are rarely caused by complex interactions of 5 or 10 parameters. Instead, they are usually: 1. **Single-mode faults:** A bug caused by a single parameter's value [24]. 2. **Double-mode faults:** A bug triggered only when a specific **interaction of two parameters** occurs (e.g., Duplex printing selected AND Printer Model 394 selected) [25-27].

 IMPORTANT

> **Pairwise Testing (All-Pairs Testing)** exploits this behavioral reality. By guaranteeing that every possible pair of parameter values is tested together at least once, we detect the vast majority of software defects (single and double-mode faults) without exhaustively testing all combinations [28-30].

Part 2: Coverage Guarantees and Reduction

- **Orthogonal Arrays / Pairwise Generation:** Test cases are carefully combined so that any given pair of inputs (e.g., Pressure \$A\$ and Temperature \$T1\$) appears in the same test case as other variables (e.g., Speed \$S1\$). We do not need a unique test case for every parameter combination, just an overlapping set [28, 31, 32].
- **Test Suite Reduction:**
 - 7 boolean parameters: Reduced from 128 exhaustive combinations to just **15 test cases** [5].
 - The reduction allows testers to achieve near 99% testing completeness with a highly optimized, physically executable number of tests [4].

Lecture 55: White box Testing

Part 1: White-Box / Structural Testing Fundamentals

White-box testing relies on the internal structure, logic, and control flow of the source code to design test cases [1], [2]. A fundamental concept in testing strategy selection is distinguishing between **strong testing** and **weak testing** [1], [3]. * **Subsumption:** If Testing Strategy 1 covers all the program elements executed by Testing Strategy 2, Strategy 1 is considered a stronger test [3], [4]. When a strong test is performed, the weaker test does not need to be explicitly executed because its coverage is automatically guaranteed [1], [3]. * **Complementary Testing:** If two strategies execute mutually exclusive or overlapping but distinct program elements, they are complementary, and both must be performed to ensure proper coverage [3], [4].

IMPORTANT

> A testing strategy is deemed "strong" only if it completely subsumes the weaker testing strategy's coverage paths [3], [5].

Part 2: Statement Coverage

Statement Coverage is the most basic structural testing metric. The objective is to design test cases such that every statement in the source code is executed at least once [2].

- **Rationale:** Bugs hidden within specific statements can only be exposed if those statements are executed during testing [2], [6].
- **Measurement:** Calculated as $\frac{\text{Number of executed statements}}{\text{Total number of statements}} * 100$ [7].
- **Limitations:** Statement coverage is a weak testing strategy [6]. Executing a statement once provides no absolute guarantee that all bugs within it or surrounding its logic will be found [6], [7]. Furthermore, in complex logic (e.g., Euclid's GCD algorithm with multiple conditional loops), finding the exact input values to achieve 100% statement coverage can be practically difficult [8], [9], [10].

Part 3: Branch Coverage / Decision Coverage

Branch Coverage (also known as **Decision Coverage**) requires that every branch of a control structure (e.g., `if`, `while`, `for`) is evaluated to both its **True** and **False** outcomes at least once [11], [12].

- **Measurement:** For n decision statements, there are $2 * n$ branches to evaluate [12], [13].
- **Subsumption Rule:** Branch coverage is strictly stronger than Statement coverage [5]. Because every statement resides within some branch, achieving 100% branch coverage mathematically guarantees 100% statement coverage [5], [14].
- **The Reverse is False:** Statement coverage does *not* guarantee branch coverage [15]. For example, in the code `if (c1) then a=b`, providing an input where `c1=true` achieves 100% statement coverage (as `a=b` executes), but completely misses testing the `c1=false` branch [16], [17].

⚠ WARNING

> While Branch Coverage is stronger than Statement Coverage, it is still insufficient for composite conditions [18]. If a decision is composite (e.g., `High = 1 OR Low = -1`), branch coverage can be achieved by satisfying just one sub-condition, masking fatal flaws in the evaluation of the other sub-conditions [18], [19], [20].

Lecture 56: Condition Testing

Part 1: The Limits of Decision Coverage

When decision statements contain multiple boolean sub-expressions (clauses) joined by logical operators (AND/OR), evaluating the entire decision to True and False is inadequate [21]. For example, if `c == alphabet OR c == digit`: * An input of an alphabet triggers the `True` branch (via short-circuiting) [22]. * An input of a special character triggers the `False` branch [23]. * Branch coverage is satisfied, but the specific logic handling `c == digit` was completely ignored and untested [24]. If a developer wrote `c = digit` instead of `c == digit`, Branch Coverage would miss the bug entirely [24].

Part 2: Basic Condition Coverage

To resolve the limitations of Branch Coverage, **Basic Condition Coverage** requires that *each individual sub-expression (clause)* within a complex decision statement must

evaluate to both **True** and **False** independently [25], [26].

- **Example:** For the decision `a > 10 AND b < 50` [26].
 - Test Case 1: `a = 15` (True), `b = 30` (True) [27].
 - Test Case 2: `a = 5` (False), `b = 60` (False) [27].
 - Both clauses have now evaluated to True and False [28].
- **Weakness:** Basic Condition Coverage does *not* guarantee Branch Coverage [29], [30]. In the expression `(a > 5 AND b < 3) OR c = 0`, test cases can be designed to flip every internal clause to T and F, yet the overall decision outcome might remain True for all test cases, meaning the False branch of the program is never executed [29], [30], [31]. Thus, Condition Coverage and Decision Coverage are incomparable; neither subsumes the other [31], [32].

Part 3: Decision/Condition Coverage

To ensure a more robust test, **Condition/Decision Coverage** combines the requirements of both prior strategies [33], [34]. * Every sub-condition must evaluate to True and False [33]. * The overall decision must evaluate to True and False [34]. * This strategy is strictly stronger than both Basic Condition Coverage and Branch Coverage [34].

Part 4: Multiple Condition Coverage (MCC)

Multiple Condition Coverage (MCC) requires testing *all possible truth value combinations* of the sub-expressions [34], [35]. * If there are n sub-expressions, MCC requires exactly 2^n test cases [35], [36]. * **The Impracticality of MCC:** While MCC is exhaustively thorough, it suffers from exponential explosion [36]. A single decision statement with 20 sub-expressions would require 2^{20} (over 1 million) test cases, making it practically impossible to implement in real-world scenarios, particularly in embedded and safety systems [36], [37], [38].

Lecture 57: MC/DC Coverage

Part 1: Modified Condition/Decision Coverage (MC/DC)

Because MCC's 2^n test case requirement is mathematically prohibitive, **Modified Condition/Decision Coverage (MC/DC)** was developed to provide the rigor of MCC but with a linear number of test cases [39], [40].

NOTE

> **Why MC/DC for Safety-Critical Software?** > Safety-critical systems require a rigorous demonstration that every single logical condition acts as intended without combinatorial explosion [37], [39]. MC/DC solves this by ensuring that every condition independently affects the decision outcome, reducing the test cases from 2^n to a much more manageable number (typically $N+1$), while still satisfying rigorous industry testing standards [40], [41], [42].

Core Definition requirements of MC/DC: 1. The overall decision takes both True and False outcomes [43]. 2. Every sub-expression (condition) takes both True and False outcomes [44]. 3. **Independence Pair Condition:** Each sub-expression must be shown to *independently affect* the outcome of the overall decision [41]. This is achieved by holding all other conditions constant and toggling the condition of interest to see the final decision toggle [45], [46].

Part 2: Finding MC/DC Test Cases (The Extended Truth Table)

To find the required test cases for an expression like `A AND B AND C`, we build an **Extended Truth Table** to map the independent effect of each variable [47], [48], [49].

MC/DC Pair Derivation Matrix Strategy: 1. Generate the standard truth table [48]. 2. For each variable (e.g., Variable A), find a pair of rows where: - Variable A toggles (1 to 0 or 0 to 1) [50]. - The outcome toggles (1 to 0 or 0 to 1) [50], [51]. - All other variables (B, C) remain strictly constant [50], [52].

Example Matrix (A AND B): | Row | A | B | Outcome | Notes for Independence Pair | | :--- | :--- | :--- | :--- | | 1 | T | T | T | A is True, B is True -> Outcome True [53] | | 2 | T | F | F |

Independent effect of B (compare Row 1 & 2) [54] || 3 | F | T | F | Independent effect of A (compare Row 1 & 3) [53] || 4 | F | F | F ||

- **To prove A's independence:** Use Row 1 and Row 3. (A toggles, B is constant at True, Outcome toggles) [53], [55].
- **To prove B's independence:** Use Row 1 and Row 2. (B toggles, A is constant at True, Outcome toggles) [54], [55].
- **Resulting Minimal Test Set:** Rows {1, 2, 3} form the MC/DC test suite [55]. Notice that for $n=2$ variables, we need $n+1 = 3$ test cases, vastly outperforming MCC as n grows [55], [42].

Coverage Hierarchy Summary Diagram

Statement Coverage < Branch/Decision Coverage < Condition/Decision Coverage < MC/DC < Multiple Condition Coverage (MCC) [56], [34], [57].

Part 3: Control Flow Graph (CFG) Fundamentals

Path testing relies on building a **Control Flow Graph (CFG)**, a directed graph visualizing the execution flow of the code [58], [59].

- **Nodes:** Represent individual statements or sequential blocks of program parts [58].
- **Edges:** Directed arrows representing the transfer of control from one node (statement) to another [58], [59].
- **Basic Statement Types in CFG:**
 1. **Sequence Statements:** Edges linearly connect one node to the next (e.g., Node 1 -> Node 2) [60].
 2. **Selection (Decision) Nodes:** A single node splits into multiple outbound edges representing branching logic (e.g., `if a>b`). Control flows to either Node 2 or Node 3 based on the evaluation [61]. These branches later merge at a **Join Node** (Node 4) to resume sequential flow [62].
 3. **Iteration (Looping) Nodes:** A `while` or `for` loop evaluates a decision. If True, control flows into the loop body. Critically, after the loop body executes, the edge must loop back to the *Decision Node* for re-evaluation [63]. Only when the decision evaluates to False does the control edge point to the next sequential statement outside the loop [63], [64].

Lecture 58: MC/DC Testing

Part 1: The Problem with Multiple Condition Coverage (MCC)

When testing complex decision statements with multiple boolean sub-expressions, the ultimate goal is often to ensure exhaustive coverage. However, **Multiple Condition Coverage (MCC)** requires generating test cases for all possible combinations of truth values. * If a boolean expression has n sub-expressions, MCC requires 2^n test cases [1]. * For example, even with just 5 sub-expressions, MCC requires 32 test cases [2]. * Because the number of test cases increases exponentially, MCC is highly impractical for real-world software, especially in embedded and control applications where expressions can have up to 20 sub-expressions [3], [4].

Part 2: Introduction to MC/DC (Modified Condition/Decision Coverage)

To overcome the exponential explosion of MCC while maintaining a highly rigorous test standard, **MC/DC (Modified Condition/Decision Coverage)** is used. MC/DC is a crucial testing criterion that enforces a rigorous level of coverage without requiring 2^n test cases [1], [5].

IMPORTANT

> **Core Principle of MC/DC:** > Every condition within a decision must take on all possible outcomes at least once, and it must be shown to **independently affect the decision's outcome** [6], [7].

MC/DC Requirements:

1. The decision must take both **True** and **False** outcomes [8].
2. Every individual sub-expression (condition) must take both **True** and **False** outcomes [8].
3. Each sub-expression must independently affect the overall decision's outcome [8].

4. To prove this independent effect, you must hold all other sub-expressions constant while varying the target sub-expression from True to False, observing that the overall decision also changes from True to False (or vice versa) [9], [10].

TIP

> MC/DC achieves complete branch coverage and is extremely powerful, yet it only requires $n + 1$ test cases (where n is the number of base conditions) [11].

Part 3: Creating an MC/DC Test Suite

The formal methodology to derive MC/DC test cases involves constructing truth tables [1], [12].

Step-by-Step Example Process:

1. **Construct a Truth Table:** Create a truth table containing all variables (e.g., A, B, C) and the final decision outcome. For 3 variables, this means $2^3 = 8$ rows [13], [14].
2. **Identify Independent Effect:** For each variable, scan the truth table to find pairs of rows where the target variable's value changes (from 1 to 0), the final outcome changes correspondingly (from 1 to 0), and all other variables remain constant [9], [10].
3. **Build the Extended Truth Table:** Mark which rows demonstrate the independent effect for each specific variable [15], [16].
4. **Select Minimum Test Suite:** Choose a minimal combination of rows (test cases) that satisfy the independent effect for *all* variables. For example, selecting rows 2, 3, 4, and 6 can cover the requirements for A, B, and C with only a fraction of the test cases needed for full MCC [17], [18].

Lecture 59: Path Testing

Part 1: Control Flow Graphs (CFG)

Path Testing is defined by modeling the execution of a program through a **Control Flow Graph (CFG)** [19].

To draw a CFG, program statements are mapped as **nodes**, and the flow of control between these statements are mapped as **edges** [19], [20]. There are three fundamental

structural constructs in any program: 1. **Sequence:** Control flows linearly from one statement to the next (Node 1 $\rightarrow$ Node 2) [21], [22]. 2. **Selection (Decision):** Statements like `if/else` where control splits into multiple branches (e.g., Node 1 branches to Node 2 or Node 3, before converging at Node 4) [22], [23]. 3. **Iteration (Loop):** Statements like `while` loops. The condition evaluates; if true, control enters the loop (Node 1 $\rightarrow$ Node 2 $\rightarrow$ Node 3) and must return to the condition node (Node 3 $\rightarrow$ Node 1) for re-evaluation. Control only proceeds to the next statement (Node 4) when the condition evaluates to false [23], [24].

Part 2: Basis Path Testing

NOTE

> **Why do we need Basis Path Testing?** > If a program contains loops, the total number of possible paths can be infinite [25]. It is impossible to test all paths. Basis Path Testing solves this by identifying a finite set of **Linearly Independent Paths**.

- **Linearly Independent Path:** A path that introduces at least one new edge that is not included in any previously defined independent path [25].
- By testing this minimal set of independent paths (the basis set), we ensure comprehensive coverage without getting trapped in infinite loop permutations [25], [26].

Part 3: Cyclomatic Complexity $V(G)$

McCabe's Cyclomatic Complexity metric evaluates the structural complexity of a program and identifies the exact number of linearly independent paths [27], [28].

This metric is a direct indicator of the psychological complexity of a program, the difficulty of testing, understanding, and debugging, as well as the number of bugs likely present [29], [30], [31]. Many organizations mandate keeping $V(G) \leq 10$ per module to maintain code reliability [32].

Calculation Methods for $V(G)$: 1. $V(G) = E - N + 2P$: Where E is the number of edges, N is the number of nodes, and P represents connected components (typically $P=1$ for a single program graph, simplifying to $E - N + 2$) [28], [29]. 2. $V(G) = \text{Number of Bounded Regions} + 1$: Bounded (closed) regions can be visually counted on the CFG, and 1 is added for the outer unbounded region [28], [29]. 3. $V(G) =$

\text{Number of Decision Nodes} + 1: Counting the number of decision/predicate statements (like `if`, `while`, `for`) in the code and adding 1 gives the exact complexity [33], [34].

Lecture 60: Dataflow and Mutation Testing

Part 1: Data Flow Testing

Data Flow Testing is a white-box testing technique focused on analyzing how variables are defined and utilized throughout a program [35]. It ensures that data anomalies (like using an uninitialized variable) do not exist.

Core Terminology:

- * **DEF (Definition):** A statement SS where a variable XX is assigned a value (e.g., `b = a` defines `b`) [36].
- * **USE:** A statement $S1$ where the value of a variable XX is read or processed (e.g., `P = X * 2` uses `X`) [36], [37].
- * **Live Variable:** A variable XX is "live" at statement $S1$ if it was defined at SS , and there exists a path from SS to $S1$ where XX is *not* redefined [38], [37].
- * **DU-Chain (Definition-Use Chain):** A set consisting of a definition of a variable and all reachable uses of that variable without any intervening redefinitions [39], [40].
- * **DU-Path:** The actual control-flow path traversed from the DEF node to the USE node [41].
- * **DC-Path (Definition-Clear Path):** A path where the variable retains its defined value without being overwritten.

Coverage Criteria:

- * **All-defs coverage:** Requires test cases to traverse at least one path from every definition of a variable to at least one of its uses.
- * **All-uses coverage:** Requires test cases to traverse at least one path from every definition of a variable to every possible use of that definition [39], [40].
- * **All-du-paths coverage:** The most rigorous criteria, requiring execution of all possible definition-clear paths from every definition to every use.

Part 2: Mutation Testing

Mutation Testing is an advanced **fault-based testing** strategy used to evaluate the effectiveness and adequacy of an existing test suite [42], [43].

 NOTE

> **Why use Mutation Testing?** > Standard coverage criteria (like statement or branch coverage) do not guarantee that bugs are actually caught, only that lines of code were executed. Mutation Testing measures how good the test suite is at actively finding injected faults [42].

Mutant Generation and Operators

The process involves making small, syntactic modifications to the original program to create **Mutant programs** [42], [43]. **Common Mutant Operators:** * Statement deletion [44]. * Replacing one statement with another [44]. * Replacing relational operators (e.g., `<` changed to `<=`) [44]. * Replacing boolean expressions (e.g., changing entirely to `True` or `False`) [44]. * Replacing arithmetic operators (e.g., changing `*` to `+`) [45]. * Variable replacement (exchanging variables of the same scope/type) [45].

Mutant Evaluation

Once millions of mutants are generated [44], they are executed against the existing test suite: * **Killed Mutants:** If the test suite produces a different output for the mutant compared to the original program, the mutant is considered "killed" (the test suite successfully caught the fault) [42], [43]. * **Live Mutants:** If the test suite produces the exact same output for the mutant as the original program, the mutant survives [42], [43]. This indicates weaknesses in the test suite, requiring the addition of new test cases. * **Equivalent Mutants:** Syntactically different but logically identical programs. No test suite can kill them.

Mutation Score Calculation Formula:
$$\text{Mutation Score} = \frac{\text{Number of Killed Mutants}}{\text{Total Number of Mutants} - \text{Number of Equivalent Mutants}}$$

Part 3: The Underlying Hypotheses of Mutation Testing

The effectiveness of Mutation Testing relies on two fundamental assumptions proposed by DeMillo (1978) [46]: 1. **Competent Programmer Hypothesis:** Assumes programmers are highly competent and write code that is very close to being correct. When they make errors, they are usually minor syntactic mistakes (e.g., swapping a `+` for a `-`) rather than

massive logical failures [45], [46]. 2. **Coupling Effect:** Posits that complex, algorithmic errors are merely combinations of several simple errors. Therefore, a test suite capable of catching simple faults (simple mutants) will inherently catch complex faults as well [47], [46].
