

EGR 430
Parallel Processing LECTURE & LAB
FINAL-EXAM(Comprehensive!)

J. Wunderlich, PhD

This is a closed-book, closed-notes, no-calculators exams. Write all answers on blank paper supplied to you. Write your name at the top of every page

BACKGROUND BASICS (XX points total)

SOURCE: Powers of 10 & 2

- | | |
|--|--|
| What power of 10 is Tera? ____ | What approx. power of 2 is Tera? ____ |
| What power of 10 is Giga? ____ | What approx. power of 2 is Giga? ____ |
| What power of 10 is Mega? ____ | What approx. power of 2 is Mega? ____ |
| What power of 10 is Kilo? ____ | What approx. power of 2 is Kilo? ____ |
| What power of 10 is Milli? ____ | |
| What power of 10 is Micro? ____ | |
| What power of 10 is Nano? ____ | |
| What power of 10 is Pico? ____ | |

SOURCES: Throughout this course, and prerequisite courses

Define these acronyms in computing:

CPU	GPU	FSB
MPP	SMP	LAN
IPC	DOP	DMA
CPI	FLOPS	
RISC	CISC	
ISA	HBM <u>Answer: High Bandwidth Memory</u>	TC <u>Tensor Core</u>
FET	MOS	TSMC (hint, it's a company)

SUPERCOMPUTER FUNDAMENTALS [PPTX](#) [PDF](#) [MP4](#) [YouTube](#)

Concerning **MULTI-CORE / MULTIPROCESSOR architectures**, sketch and fully label a picture including Cores/Processors, Cache's, Memory, and interconnectivity of the five different architectures discussed in lecture:

1. simple single processor
2. SMP
3. Vector-Register
4. MPP
5. large network dedicated to a single task

• SOURCES: SLIDE 42 of **SUPERCOMPUTER FUNDAMENTALS** [PPTX](#) [PDF](#) [MP4](#) [YouTube](#)

Assuming a simple 4-stage machine Instruction execution cycle (Fetch, Decode, Execute, Wright Back) **in a single Core or Processor**, make three graphs (with instructions on the Y axis and time on the X axis) showing the difference between two instructions, executing on

1. (graph-1) A **NON-PIPELINED** processor or core
2. (graph-2) a **PIPELINED** processor or core, and
3. (graph-3) A **TWO-WAY** (i.e., "two-issue") **SUPERSCALAR** processor or core.

PARALLEL PROCESSING HARDWARE & SYSTEMS-LEVEL SOFTWARE OPTIMIZATION [PPTX](#) [w/audio](#) [PDF](#) [MP4](#) [YouTube](#)

- State a type of **DATA DEPENDENCY** in hardware
 - (ANSWER: when the execution of a machine instruction in one of the parallel pipelines in a superscalar architecture is dependent on another machine instruction in one of the parallel pipelines)
- State a type of **SOURCE DEPENDENCY** given in lecture
- State a TWO types of **CONTROL DEPENDENCIES** given in lecture

PARALLEL INTERCONNECT ARCHITECTURES of Cores, Processors, and Networks [PPTX](#) [w/audio](#) [PDF](#) [MP4](#) [YouTube](#)

- A. For a ____-node ____ topology LAN (static interconnect network architecture) with serial-packetized TCP/IP communication interconnection links, calculate:
- i. Number of Nodes, AND say how to optimize
 - ii. Network "Diameter", AND say how to optimize
 - iii. Node "Degree", AND say how to optimize
 - iv. Network "Bisection Width", AND say how to optimize (Max and Min)
 - v. The total number of physical wires through the Network Bisection, AND say how to optimize
- B. For a ____-node ____ topology MPP (static interconnect network architecture) connected with 64-bit parallel links between each node, calculate:
- i. Network "Diameter", AND say how to optimize
 - ii. Node "Degree", AND say how to optimize
 - iii. Network "Bisection Width", AND say how to optimize
 - iv. The total number of physical wires through the Network Bisection, AND say how to optimize (Max and Min)
- C. Sketch two different **dynamic** interconnect network architectures

PARALLEL PROCESSING HARDWARE & SYSTEMS-LEVEL SOFTWARE OPTIMIZATION [PPTX](#) [w/audio](#) [PDF](#) [MP4](#) [YouTube](#)

— SOURCES: [PPTX](#) [PDF](#) [MP4](#) [YouTube](#) AND [PPTX](#) [w/audio](#) [PDF](#) [MP4](#) [YouTube](#)

Given the following equations of SCALAR equation: _____

1. Graph instructions (in circles) and arrows indicating flow of execution, for a SOFTWARE (IDEAL) implementation, AND calculate the **AVERAGE PARALLELISM (DOP)** = (# of instructions) / (# of cycles)
2. Graph instructions (in circles) and arrows indicating flow of execution, for a ____-WAY SUPERSCALAR PROCESSOR HARDWARE implementation, AND calculate the **AVERAGE PARALLELISM (DOP)** = (# of instructions) / (# of cycles). ASSUME:
 - memory accesses ("load" or "store") can be done simultaneously
 - arithmetic instruction can be done at a time in this one core/processor
 - You must use STORE instructions in the last cycle(s) of your execution
3. Graph instructions (in circles) and arrows indicating flow of execution, for ____ CORE SMP HARDWARE (communicating through "Loads" & "Stores"), AND calculate the **AVERAGE PARALLELISM** = (# of instructions) / (# of cycles). ASSUME:
 - i. **ONLY 1** memory access ("Load" or "Store") can be done at a time on **EACH** core/processor
 - ii. **OR 1** arithmetic instruction can be executed at the same time as a Load/Store on **EACH** core/processor
 - iii. **i.e., they are single-issue non-superscalar processors**

Possibly do this exact problem on Test:

GIVEN: $C = \vec{a} \cdot \vec{b}^T = \begin{bmatrix} a_{11} & a_{12} \end{bmatrix} * \begin{bmatrix} b_{11} \\ b_{21} \end{bmatrix}$

VECTOR DOT-PRODUCT → $C = (a_{11} * b_{11}) + (a_{12} * b_{21})$

SMP ARCHITECTURE

Diagram: Two processors (P) connected to a shared memory (M).

Definitions:

- SMP = SYMMETRIC MULTI-PROCESSING
- M = MAIN MEMORIES (DRAM)
- DMA = DIRECT MEMORY ACCESS
- IPL = INTER-PROCESSOR COMMUNICATION
- DOP = DEGREE OF PARALLISM (AVE. HARDWARE PARALLELISM)
- SPEEDUP = TIME FOR WORST CASE / TIME FOR NEW DESIGN

Annotations:

- * TAKES 8 CLOCKS
- + TAKES 2 CLOCKS
- L LOADS/STORES TAKES 5 CLOCKS
- T TAKES 10 CLOCKS

A) MAKE FINE-GRAIN GRAPH AND FIND IDEAL(AUSSERWAKE) DOP

B) SCHEDULE THIS ON A SINGLE-ISSUE P

C) CALCULATE DOP

Diagram: A graph showing nodes L1, L2, L3, L4, A, B, J, and SUM. Edges represent dependencies. The graph is annotated with "CYCLES" and "INSTRUCTIONS".

Diagram: A timeline diagram showing operation cycles for L1, L2, L3, L4, A, B, J, and SUM. The timeline is labeled "OPERATION CYCLE #".

Timeline Data:

Operation Cycle #	Task
0	L1
5	L1
10	L2
15	L2
20	L3
25	L3
30	L4
35	L4
40	A
45	A
50	B
55	B
60	J
65	J
70	SUM
75	SUM

Annotations:

- DOP = $\frac{8}{8} = 1$ WORST POSSIBLE DOP
- MEAN
- WORST CASE TIME

D) SCHEDULE SAME FINE-GRAIN PROGRAM GRAPH ON A 4-WAY SUPERSCALAR PROCESSOR

Diagram: A table showing operation cycles for L1, L2, L3, L4, A, B, J, and SUM.

Table Data:

Operation Cycle	L1	L2	L3	L4
1				
2	5			
3	15			
4	23			
5	28			
6		15		
7		23		
8		28		
9			15	
10			23	
11			28	
12				15
13				23
14				28
15				
16				
17				
18				
19				
20				
21				
22				
23				
24				
25				
26				
27				
28				
29				
30				
31				
32				
33				
34				
35				
36				
37				
38				
39				
40				
41				
42				
43				
44				
45				
46				
47				
48				
49				
50				
51				
52				
53				
54				
55				
56				
57				
58				
59				
60				
61				
62				
63				
64				
65				
66				
67				
68				
69				
70				
71				
72				
73				
74				
75				
76				
77				
78				
79				
80				
81				
82				
83				
84				
85				
86				
87				
88				
89				
90				
91				
92				
93				
94				
95				
96				
97				
98				
99				
100				

E) DOP = $\frac{8}{4} = 2$

F) SPEEDUP = $\frac{53}{28} = 1.96$

Annotation: 96% FASTER THAN WORST CASE

G) OBSERVE WASTED RESOURCES IN 2-WAY - PROCESSOR DESIGN

H) AND USE COURSE-GRAIN PACKING TO OPTIMIZE AND IMPLEMENT THEM ON A Z-PROCESSOR SYSTEM

I) MAKE A TABLE COMPARING DOP'S AND SPEED-UPS

J) DRAW CONCLUSIONS

K) ALTHOUGH DOP IS GOOD FOR CONCEPTUAL DESIGN OF PARALLEL ARCHITECTURES, WE:

- NEED TO SCHEDULE ACTUAL TIMES, AND THEN CALCULATE SPEED-UP
- CONSIDER WASTED RESOURCES, CONSIDER A MORE EFFICIENT ARCHITECTURE AND DO GRAIN-PACKING, TO MAYBE OPTIMIZE SPEED-UP

L) EVEN THOUGH INTENTIONALLY SMALL IPC

M) NOT SO FOR THIS EXAMPLE

N) WORST CASE

O) 25% FASTER THAN WORST CASE

P) 7

Q) 10

R) 1.43

S) 53

T) 43%

U) 1.23

V) 25%

W) FASTER

X) THAN

Y) WORST

Z) CASE

A) OPERATION TOTAL

B) TIME

C) L

D) S

E) V

F) W

G) S_i

H) L_s

I) Z

J) P_r

K) P_s

L) Intentionally small to make a point

M) DOP (AVE. HARDWARE PARALLELISM)

N) SPEED-UP

O) IDEAL SINGLE PROCESSOR

P) 2 BEST

Q) 1 WORST

R) 2 BEST

S) 1.43

T) N.A.

U) DATUM

V) 96% BEST

W) 23%

X) TWO-SWAY SUPERSCALAR

Y) TWO SINGLE-ISSUE PROCESSORS

Z) AND GRAIN PACKING

3) For this problem, grain packing on a multi-core system doesn't make a difference since it's such a small computation, best done by a single Superscalar Processor; which is 4-Way here, however even if it was only 2-Way, with therefore $DOP = 8/5 = 1.6$, and $Speedup = 53/33 = 1.6$ (i.e., 60% faster than the worst case), superscalar is still better than the 2-processor case

I May ask you a piece of this problem:

PARALLEL PROCESSING HARDWARE & SYSTEMS-LEVEL SOFTWARE OPTIMIZATION PPTX-w/audio PDF MP4 YouTube

$$\begin{bmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{bmatrix} \times \begin{bmatrix} B_{11} & B_{12} \\ B_{21} & B_{22} \end{bmatrix} = \begin{bmatrix} C_{11} & C_{12} \\ C_{21} & C_{22} \end{bmatrix}$$
$$SUM = C_{11} + C_{12} + C_{21} + C_{22}$$

SOURCES: PPTX PDF MP4 YouTube AND PPTX-w/audio PDF MP4 YouTube

A. Given the following **MATRIX** (i.e. Vectors), and **SCALAR** equations

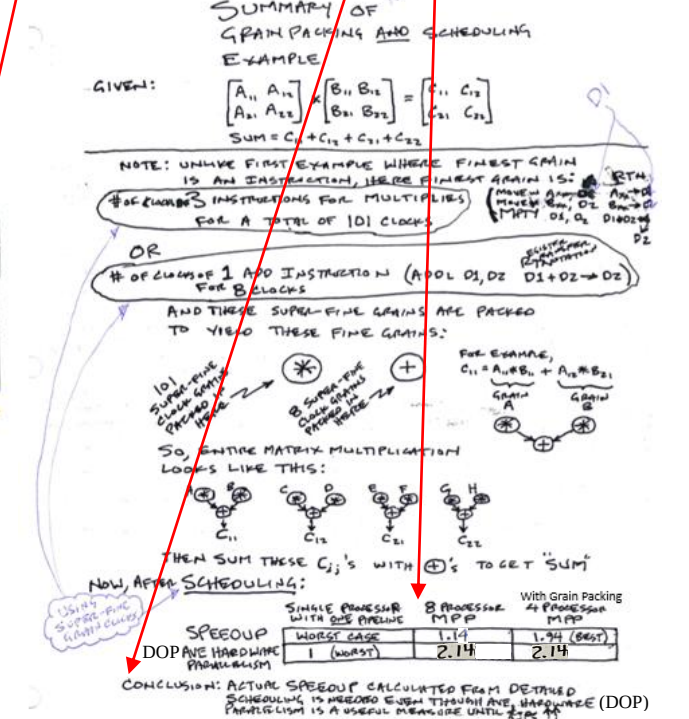
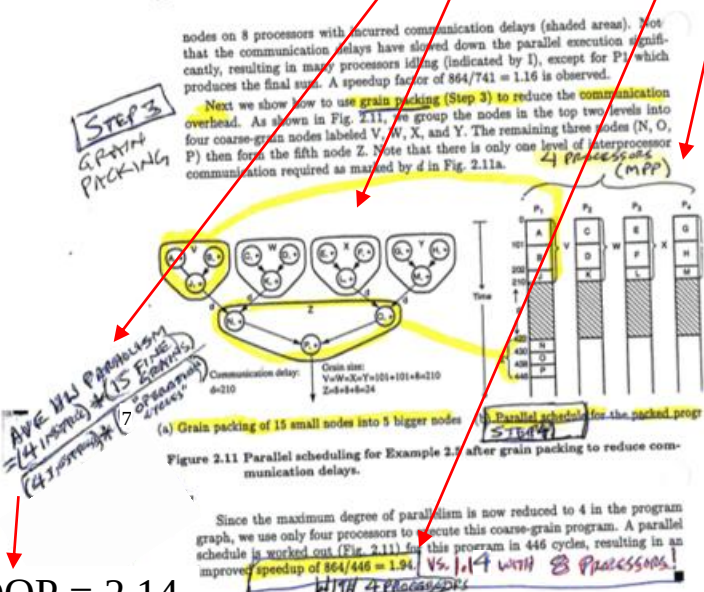
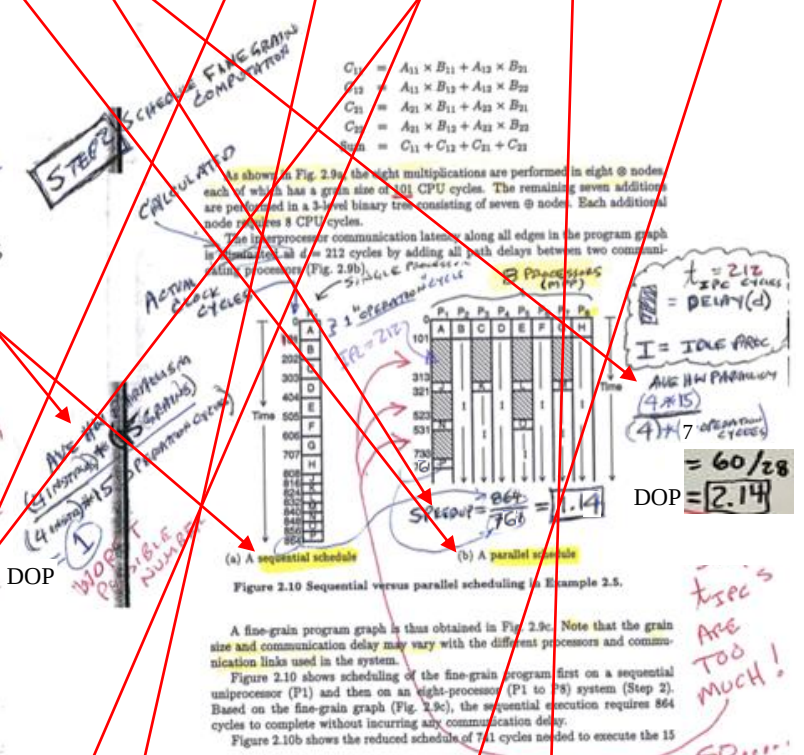
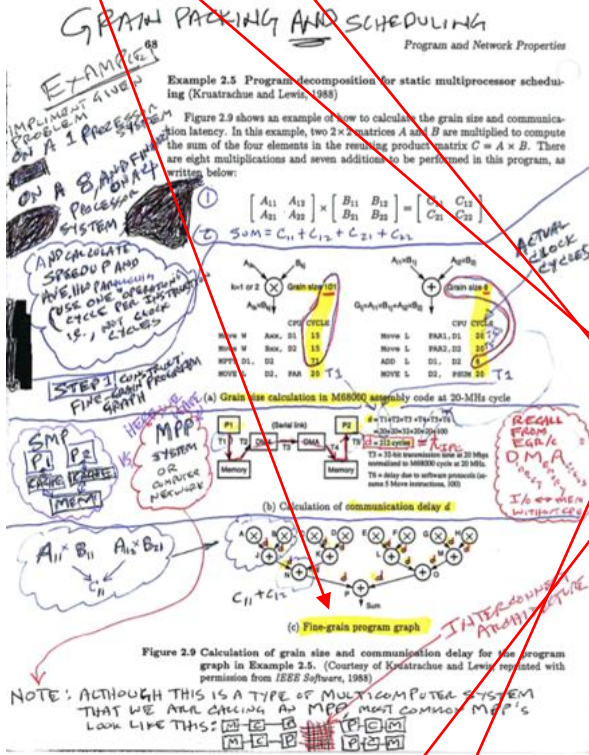
Compare (with **TABLE** and **CONCLUSION**), **SHOWING ALL CALCULATIONS** for **AVERAGE HARDWARE PARALLELISM (DOP)**, and **SPEEDUP** (**SHOW SCHEDULES WITH EXACT TIMINGS**) on:

- A single processor (non-pipelined) (**DRAW FINE-GRAIN GRAPH**)
- An 8-core MPP (non-pipelined cores), using DMA's for IPC (**USING SAME FINE-GRAIN GRAPH**)
- A 4-core MPP (non-pipelined cores), using DMA's for IPC, and **GRAIN-PACKING** (**DRAW COURSE-GRAIN GRAPH**)

Assume 4 Fine-Grain Assembly Instructions for **MULTIPLICATION-Grains** (needing **101 "Clocks"**), and **ADD-Grains** (needing **8 "Clocks"**), and **IPC = 212 "clocks"** ... For **DOP**, assume one operation-Cycle per Fine-Grain (i.e., not "Clocks"), and also one operation-Cycle per Course-Grains after Grain-packing

HINT: draw this, draw this, calculate this,

draw this, calculate these, draw this, draw this, calculate these, make this table, state this conclusion



Same problem as above:

GIVEN: $A * B = C$

MATRICES

$A = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix}$ $B = \begin{bmatrix} b_{11} & b_{12} \\ b_{21} & b_{22} \end{bmatrix}$ $C = \begin{bmatrix} c_{11} & c_{12} \\ c_{21} & c_{22} \end{bmatrix}$ AND $SUM = C_{11} + C_{12} + C_{21} + C_{22}$

VECTOR DOT-PRODUCT OF ROW 1 OF A AND COLUMN 1 OF B

ALSO GIVEN:

$C_{11} = (a_{11} * b_{11}) + (a_{12} * b_{21})$

FOUR ASSEMBLY INSTRUCTIONS TO DO $*$

FOUR ASSEMBLY INSTRUCTIONS TO DO $+$

TAKES 101 CLOCKS **TAKES 8 CLOCKS**

MPP ARCHITECTURE WITH DMA'S ALLOWING DEDICATED MEMORIES TO TALK TO EACH OTHER WITHOUT TALKING UP PROCESSORS

MPP = MASSIVELY PARALLEL PROCESSORS
M = MAIN "MEMORIES" (DRAM)
DMA = DIRECT MEMORY ACCESS
IPC = INTER-PROCESSOR COMMUNICATION
DOP = DEGREE OF PARALLELISM (AVE. HARDWARE PARALLELISM)
SPEED-UP = TIME FOR WORST CASE / TIME FOR NEW DESIGN

1) OBSERVE 8-PROCESSOR DESIGN (ESPECIALLY P_1, P_2, P_3, P_4) AND USE COURSE-GRAIN PACKING TO OPTIMIZE AND IMPLEMENT THEM ON A 4-PROCESSOR SYSTEM

A) MAKE A FINE-GRAIN PROGRAM GRAPH

B) SCHEDULE THIS ON A SINGLE NON-SUPERSCALAR PROCESSOR

C) CALCULATE DOP

OPERATION CYCLE #

TIME

Worst Case Time

FINE GRAINS 101 CLOCKS EACH

FINE GRAINS 8 CLOCKS EACH

DOP = $\frac{(4 \text{ INSTRUCTIONS}) * (15 \text{ FINE GRAINS})}{(4 \text{ INSTRUCTIONS}) * (15 \text{ OPERATION CLOCKS})}$

= 1 WORST POSSIBLE DOP

D) SCHEDULE SAME FINE-GRAIN PROGRAM GRAPH ON AN 8-PROCESSOR MPP (ALL NON-SUPERSCALAR PROCESSORS) WITH DMA'S AND SAME TIMINGS

OPERATION CYCLE #

TIME

Worst Case Time

FINE GRAINS 101 CLOCKS EACH

FINE GRAINS 8 CLOCKS EACH

DOP = $\frac{(4 \text{ INSTRUCTIONS}) * (15 \text{ FINE GRAINS})}{(4 \text{ INSTRUCTIONS}) * (7 \text{ OPERATION CLOCKS})}$

= 60/28 = 2.14 BETTER THAN 1

F) CALCULATE SPEEDUP

864 / 761 = 1.34

1.34 FASTER THAN WORST CASE

3) OBSERVE 8-PROCESSOR DESIGN (ESPECIALLY P_1, P_2, P_3, P_4) AND USE COURSE-GRAIN PACKING TO OPTIMIZE AND IMPLEMENT THEM ON A 4-PROCESSOR SYSTEM

1) OBSERVE 8-PROCESSOR DESIGN (ESPECIALLY P_1, P_2, P_3, P_4) AND USE COURSE-GRAIN PACKING TO OPTIMIZE AND IMPLEMENT THEM ON A 4-PROCESSOR SYSTEM

2) SCHEDULE THIS ON A SINGLE NON-SUPERSCALAR PROCESSOR

3) CALCULATE DOP

OPERATION CYCLE #

TIME

Worst Case Time

FINE GRAINS 101 CLOCKS EACH

FINE GRAINS 8 CLOCKS EACH

DOP = $\frac{(4 \text{ INSTRUCTIONS}) * (15 \text{ FINE GRAINS})}{(4 \text{ INSTRUCTIONS}) * (7 \text{ OPERATION CLOCKS})}$

= 60/28 = 2.14 BETTER THAN 1

F) CALCULATE SPEEDUP

864 / 446 = 1.94

1.94 FASTER THAN WORST CASE

1) OBSERVE 8-PROCESSOR DESIGN (ESPECIALLY P_1, P_2, P_3, P_4) AND USE COURSE-GRAIN PACKING TO OPTIMIZE AND IMPLEMENT THEM ON A 4-PROCESSOR SYSTEM

2) SCHEDULE THIS ON A SINGLE NON-SUPERSCALAR PROCESSOR

3) CALCULATE DOP

OPERATION CYCLE #

TIME

Worst Case Time

FINE GRAINS 101 CLOCKS EACH

FINE GRAINS 8 CLOCKS EACH

DOP = $\frac{(4 \text{ INSTRUCTIONS}) * (15 \text{ FINE GRAINS})}{(4 \text{ INSTRUCTIONS}) * (7 \text{ OPERATION CLOCKS})}$

= 60/28 = 2.14 BETTER THAN 1

F) CALCULATE SPEEDUP

864 / 446 = 1.94

1.94 FASTER THAN WORST CASE

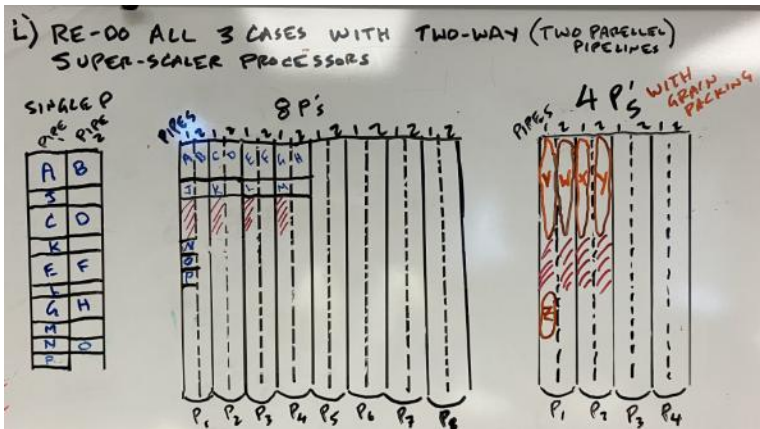
J) MAKE A TABLE COMPARING DOP'S AND SPEED-UPS FOR ALL THREE CASES

	8 PROCESSORS	4 PROCESSORS
DOP (AVE. HARDWARE PARALLELISM)	2.14	2.14
SPEED-UP	1.34	1.94 BEST

K) DRAW CONCLUSIONS

ALTHOUGH DOP IS GOOD FOR CONCEPTUAL DESIGN OF PARALLEL ARCHITECTURES, WE:

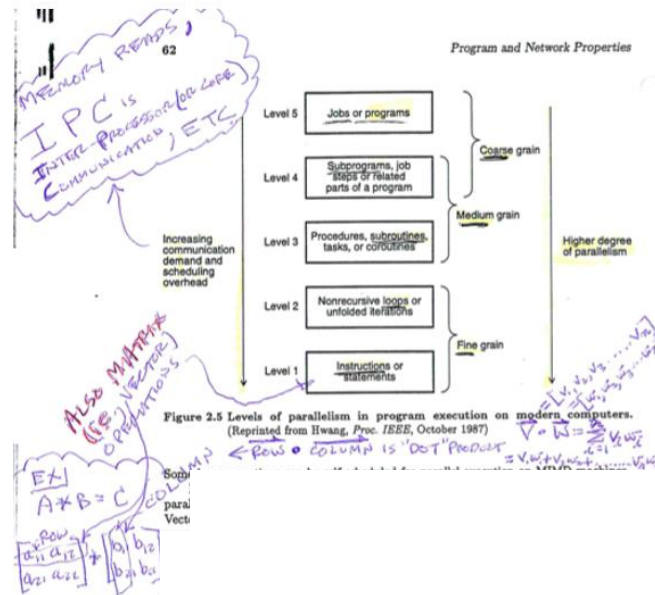
- NEED TO SCHEDULE ACTUAL TIMES, AND THEN CALCULATE SPEED-UP
- CONSIDER WASTED RESOURCES, CONSIDER A MORE EFFICIENT ARCHITECTURE AND DO GRAIN-PACKING, TO OPTIMIZE SPEED-UP



SOURCES: **PARALLEL PROCESSING HARDWARE & SYSTEMS-LEVEL SOFTWARE OPTIMIZATION** PPTX-w/audio PDF MP4 [YouTube](#)

Create a 5 tiered picture showing the five er: levels of parallelism in program execution, from level 1 and 2 being fine grained, Level 3 and 4 being medium grain, and level 5 being coarse grain... fill in the description of each of the levels include an arrow going from the top to the bottom saying something about IPC; additionally make a sketch of a matrix multiplication row times column dot product next to the level 1

Answer:



ASSEMBLY LANGUAGE

SOURCE: https://users.etown.edu/w/wunderjt/MicrocontrollerPaper_Highlighted%203.pdf

Describe the difference between an **OP-CODE** and an **OPERAND**

SOURCE: Review EGR330 **CPU Design & Assembly Language intro**

List the four basic stages of instruction pipeline (machine instruction cycle) and how they differ for **REGISTER-REFERENCE** versus **MEMORY-REFERENCE** instructions in most microprocessors and large scale systems:

SOURCE: **Wunderlich, J.T. (1999). "Focusing on the blurry distinction between microprocessors and microcontrollers." In Proceedings of 1999 ASEE Annual Conference & Exposition, Charlotte, NC: (session 3547), [CD-ROM]. ASEE Publications. . PDF DOC**

- Write the equation for the **TIME (T) to EXECUTE MACHINE LANGUAGE** and describe each variable
- What is the difference between a **PRINCETON (Von Neumann) ARCHITECTURE**, and a **HARVARD ARCHITECTURE** in microprocessor and microcontroller design? And why are the different?
- Why are **MEMORY-REFERENCE** machine instructions slower than **REGISTER-REFERENCE** most of the time for microprocessor applications, and why can this sometimes not be true in microcontrollers.
- What is the function of a **PROGRAM COUNTER (PC)** in computer architectures, AND Why is it a counter?
- How is a **STACK** used in assembly language for program calls and interrupt service routines in assembly language?
- List **ALL SIX** methods mentioned for **REDUCING the AVERAGE CYCLES PER INSTRUCTION (CPI)**
- List any two of the differences between a **MICROCONTROLLER vs. a MICROPROCESSOR**

SOURCE: **PSW**

Draw and describe all of the contents of the **8051 microcontroller STATUS REGISTER (Program Status Word)**, AND explain how two of the bits are used in controlling Registers-Banks

SOURCE: **Overview/architecture**

Describe an **EXAMPLE the MULTIPLEXING of PINS** in the **8051 microcontroller**

SOURCE: **PPT-PDF, Book-PDF**

Describe the different **ADDRESSING MODES** in the **8051 microcontroller**

SOURCES: Number Representations (PDF,PPTX-w/audio,MP4,YouTube) AND IEEE Binary Floating Point (PDF,PPTX-w/audio,MP4,YouTube)

Show all the bits to represent decimal 0.75 in IEEE BFP (Binary Floating Point) Single Precision including **sign bit, exponent field, and mantissa**. Show all the steps of your work!

SOURCES: SLIDE #8 of PDF PPTX-w/audio MP4 YouTube

How is **PCIe** different than **PCI** in the way bits are handled in digital communications; and why was this a major paradigm shift in simple computers a couple decades ago?

NVIDIA ... SOURCES: PARALLEL PROCESSING HARDWARE & SYSTEMS-LEVEL SOFTWARE OPTIMIZATION [PPTX-w/audio PDF MP4 YouTube](#)

AND **CACHE DESIGN** [PPTX-w/audio PDF MP4 YouTube](#)

Write a very short essay (one page max) about what we discussed regarding the 2022 and 2025 export controls of **NVIDIA** chips including creating THREE graphs showing the **H800, H100, and H20 processors** with **COMPUTATIONAL PERFORMANCE (FLOP's)** on the Y axis and **INTERCONNECT BANDWIDTH** on the X axis; clearly shade an area on each graph showing the extent of the export controls for both 2022 and 2025. Include in your essay how creative deep learning distillation and optimization techniques were used as a workaround to use certain processors to outperform more costly REASONING models in 2025.

HINTS:

GPU's for Deep Learning: NVIDIA H100, H800, H20

January 2025 announcement by Chinese Company "DeepSeek" that it can perform Deep Learning at a very small fraction of the costs

AI Overview

DeepSeek's ability to build powerful AI systems inexpensively stems from its innovative training techniques and efficient model design. They achieved this by focusing on training only the most crucial parts of the model, using smart memory compression, and making efficient use of hardware resources, according to Analytics Vidhya. Specifically, they employed a "mixture-of-experts" approach, which allows the model to specialize in different tasks and only activate the necessary submodels, as explained in a paper by Live Science and The New York Times.

Here's a more detailed breakdown:

Mixture of Experts (MoE):

Instead of training one massive model, DeepSeek uses multiple smaller models, each specializing in a different task. This allows them to activate only the necessary submodels for a given task, reducing computation.

Efficient Memory Compression:

They use techniques to store the model more efficiently without sacrificing performance, further lowering the cost of training and deployment.

Smart Hardware Utilization:

DeepSeek's engineering focused on maximizing the performance of available hardware instead of relying on the most expensive and powerful chips.

Multi-Token Prediction:

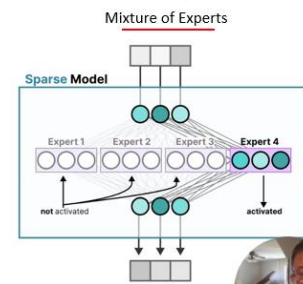
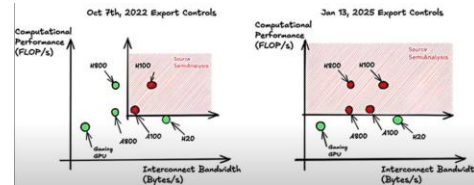
DeepSeek trained its base model to predict multiple words at once, which is cheaper and can improve accuracy, according to a report by MIT Technology Review.

Reinforcement Learning and Distillation:

DeepSeek also used reinforcement learning and distillation techniques to make their models more efficient and cost-effective, as noted on YouTube.

Open-Source Approach:

DeepSeek's open-source approach and the fact that their methods are not proprietary have also contributed to the cost-effectiveness of their models, according to Entrepreneur and a post on Threads.



Did DeepSeek Copy Off Of OpenAI? And What Is Distillation?

By John Werner, Contributor. I am an MIT Senior Fellow & Lecture, Sr...

Follow Author

Jan 30, 2025, 11:27am EST

"Distillation is a technique designed to transfer knowledge of a large pre-trained model (the "teacher") into a smaller model (the "student"), enabling the student model to achieve comparable performance to the teacher model," write Vishal Yadav and Nikhil Pandey. "This technique allows users to leverage the high quality of larger LLMs, while reducing inference costs in a production environment, thanks to the smaller student model."

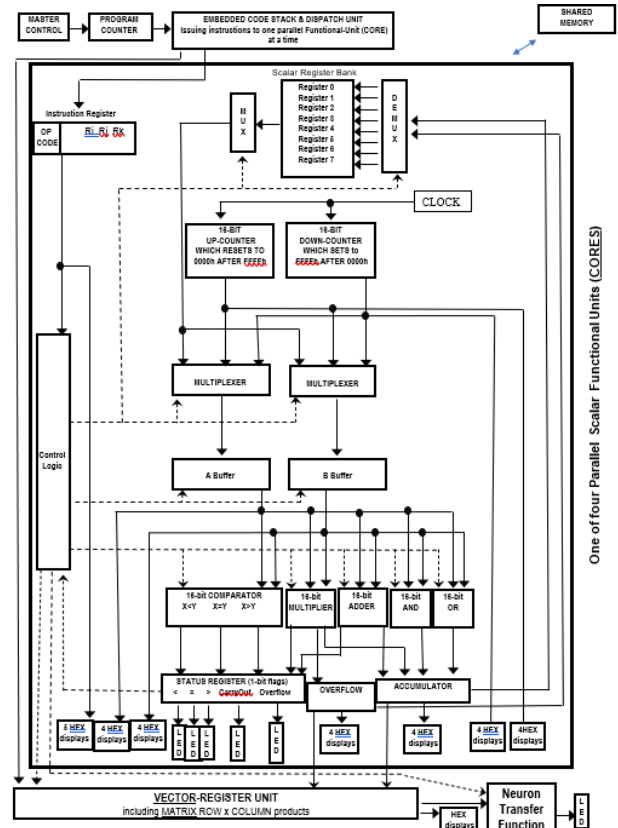
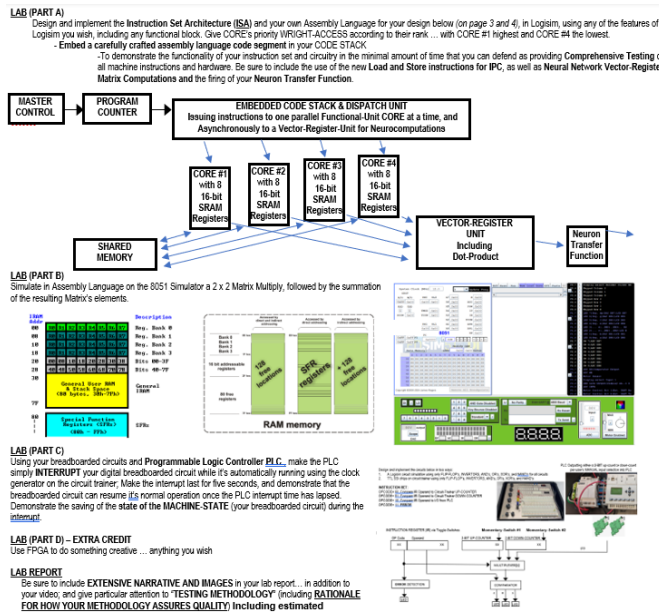
CACHE DESIGN

1. Compare and contrast **SRAM** and **DRAM**
2. define both **TEMPORAL** and **SPATIAL LOCALITY** of **REFERENCE**
3. Explain what **THRASHING** is if using direct mapping in an SMP architecture
4. Explain why a cache line replacement algorithm is not needed for a cache miss if using direct mapping in an SMP architecture
5. How is a GPU cache different than a CPU cache?
6. State both the hit and miss operation including a cache-miss replacement policy
7. For a **Main Memory** that has byte addressing, and a cache that holds of **blocks**:
 - a. For a **MAPPED** cache:
 - i. Sketch and label the cache including cache-lines (for the cache blocks) in **HEX**
 - ii. Show the calculation of the cache tag size
 - iii. Re-sketch and label Main Memory to show how it is organized ("chunked") into groups because of the size of the tag (if Direct or Set-Associative Mapping, otherwise explain why that is not necessary – i.e., if fully associative mapping is asked for)

QUALITY CONTROL / VERIFICATION PDF PPT

Reflecting on Wunderlich, J.T. (2003). Functional verification of SMP, MPP, and vector-register supercomputers through controlled randomness. In *Proceedings of IEEE SoutheastCon, Ocho Rios, Jamaica*, M. Curtis (Ed.): (pp. 117-122). IEEE Press. [PPTX-w/audio PDF MP4 YouTube](#),

- A. Explain in detail your selection of your designed Machine Instructions (Scalar, Vector, And Neural Network) that you picked for **PART A** of your last lab project below, and explain how they sufficiently tested the critical parts of the functionality of your design, in a **minimal** amount of TIME, and use of code stack memory SPACE



Write an essay about NEURAL NETWORK PROCESSOR DESIGN (Vector-Register Parallel Processing) including:

- In a sentence or two, compare and contrast the two Neurocomputer chips developed by Dr W in the early 1990's, in his draft book chapter
 - J.T. Wunderlich, Deep Learning Book Chapter Draft & Neural Network Processor Designs [PPTX-w/audio PDF MP4 YouTube](#)
- Draw and label graphs of three Neuron Transfer functions, including the Rectified Linear Unit in the 2018 paper we discussed
 - "Deep Learning using Rectified Linear Units (ReLU)" by A. Agarap - Neural Network Transfer Function Hardware
- Explain Dr W's Transfer function implementation in his single-chip parallel processing Neurocomputer with on-chip learning
- With words and sketches, describe your neural network hardware (ISA, vector register unit, and neural network transfer function unit) in your last major laboratory project, and reference some of:
 - Wunderlich Intro to Neural Networks; Data, Training, Inference [PPT PDF MP4 YouTube](#)
 - Abdi, H., Valentin, D., Edelman, B. (1999), *Neural Networks*, Sage Publications. [Deep-Learning Transformers](#)
 - Hertz, J., Krogh, A., Palmer, R. G. (1991), *Introduction to the Theory of Neural Computation*, Addison-Wesley Publishing, 1991. CH [0,1,2,3,4,5,6,7,8,9](#)
 - J.T. Wunderlich, Deep Learning Book Chapter Draft & Neural Network Processor Designs [PPTX-w/audio PDF MP4 YouTube](#)
 - "Deep Learning using Rectified Linear Units (ReLU)" by A. Agarap - Neural Network Transfer Function Hardware
 - 2024 "Benchmarking and Dissecting the Nvidia Hopper GPU Architecture" by Weile Luo, et al.

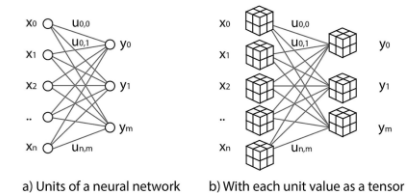
SOURCE: 2024 "Benchmarking and Dissecting the Nvidia Hopper GPU Architecture" by Weile Luo, et al.

What is a Tensor in Deep Learning?

Answer: a multidimensional array of matrices and vectors, mostly containing interneuron connection-WEIGHTS

What is a Tensor Core?

Answer: A functional unit within a GPU designed to accelerate matrix (i.e., Vector) operations



https://en.wikipedia.org/wiki/Tensor_%28machine_learning%29#/media/File:Tensor_Units.jpg

CISC vs. RISC INSTRUCTION SET DESIGN, PERFORMANCE & SCALABILITY PPTX-w/audio PDF MP4 YouTube

- SOURCE: "Simultac Fonton: A Fine-Grain Architecture for Extreme Performance beyond Moore's Law" by M. Brodowicz and T. Sterling
Explain how you believe the "Fontons" in this supercomputer paper we discussed could be used for neurocomputations
- Explain in words **AMDAHL'S LAW** as applied to multiple processors or cores in an SMP Architecture.
- Draw a graph of **AMDAHL'S LAW** as applied to multiple processors or cores in an SMP Architecture, AND also graph the ideal case.
- From reading & lecture on "BREAKING THE MULTICORE BOTTLENECK", explain in words, AND graph, how it relates to Amdahl's Law
- Other than the number and complexity of machine instructions, state 3 differences between **RISC** and **CISC** Give an example of when you may have given up too soon on trying to solve a Lab problem. Answer: **RISC** has **A hardwired control unit** which makes it faster than the micro programmed cisc control unit, **RISC** has **more general purpose registers**, the instruction format in **RISC** has **fixed boundaries between the Op code and operand fields in the Instruction Register** which make it faster to process fetched instructions

High Level Language vs Assembly Language 2 [PDF PPTX-w/audio](#) ([MP4* YouTube*](#)) –videos include #18 above (List 3 advantages and 3 disadvantages of assembly language)

SOURCE: [PDF2](#)

What is the difference between a **TRAP** and an **INTERRUPT**

SOURCES:

- Why is **VIRTUAL ADDRESSING** needed to accommodate all the expectations of software development; i.e., explain the difference between hardware realities and what we like to do when we program. Show explicitly with powers of two and magnitudes the difference between a virtual address and a real address ... i.e., Why have there only been 40 pins for **addressing memory** coming out of the back of Pentium processors even though 64 bit computing has been around for three decades.

SOURCES: [Atoms & Transistors](#)

Draw the circuit symbols and label them for both **CMOS** and **Bipolar Transistors**, then write a sentence or two about how each of them functions differently

SOURCES: [Moore's Law](#)

State Moore's law and why it must inevitably fail

SOURCES:

What is the approximate **MINIMUM FEATURE SIZE** that chips are being manufactured at this year, give a number and a Greek symbol for the dimension; then express it as a multiple of **atomic diameters (angstroms)**

Summarize in one paragraph [U.S. AMBASSADOR JOHN B. CRAIG](#), lecture on **DESIGN AND DEVELOPMENT OF COMMERCIAL AIRCRAFT**; from his experience as a Boeing Vicepresident.

Describe your favorite team moment where everybody in your group was equally happy about solving some problem together. Describe the problem, why it was so puzzling initially, and what made the solution so happy.

Recall this lab project, and sketch how you:

(a) Designed your **MULTIPLEXER** (clearly labeled block diagrams, and some narrative of what you needed to do to come up with the logic gates needed ... like cascaded bit-slices, two smaller MUX's cascaded, or maybe a method you used to simplify an 8 variable digital function.

(b) Designed your **PLC flow-chart (s)**

Design and implement the circuits below in two ways:

- A Logisim circuit simulation using only FLIP-FLOP's, INVERTORS, AND's, OR's, XOR's, and NAND's for all circuits
- TTL SSI chips on circuit trainer using only FLIP-FLOP's, INVERTORS, AND's, OR's, XOR's, and NAND's

INSTRUCTION SET:

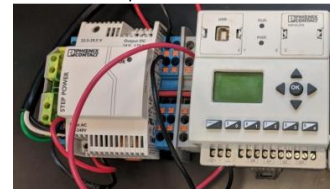
OPCODE= 00 Compare **IR** Operand to Circuit-Trainer UP-COUNTER

OPCODE= 01 Compare **IR** Operand to Circuit-Trainer DOWN-COUNTER

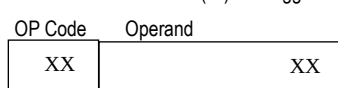
OPCODE= 10 Compare **IR** Operand to I/O from PLC

OPCODE= 11 **ERROR**

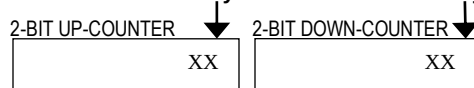
PLC Outputting either a **2-BIT** up-count or down-count per user's **MANUAL** input selection into PLC



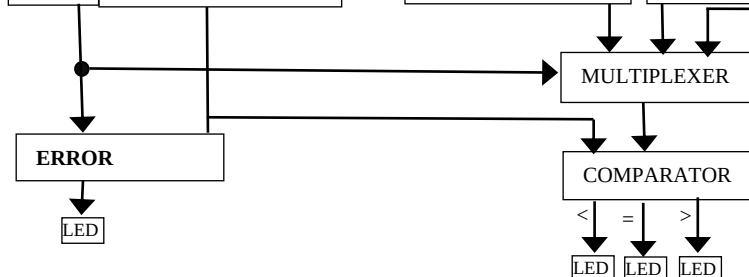
INSTRUCTION REGISTER (IR) via Toggle-Switches



Momentary-Switch #1 Momentary-Switch #2



I/O



Given the assembly language routine below,

A. With the register transfer notation blanked out on your exam, fill in the Register Transfer Notation (RTN)

B. Fill in the Memory Map with what it looks like at the end of a simulator program run (Just Write on Top of the "00"s).

; Intel 8051 Assembly Language Code to generate eleven numbers in a

; Fibonacci Series (00,01,01,02,03,05,08, 21,34,55,89) in Decimal

; Fibonacci Series (00,01,01,02,03,05,08,0D,15,22,37,59) in Hex

```

MOV R0,#30h ;put ADDRESS 30hex of MEMORY location to store Series-Number, into REGISTER R0
MOV @R0,#00h ;put DATA 01hex into MEMORY location pointed to by the address in REGISTER R0
MOV R1,#31h ;put ADDRESS 31hex of MEMORY location to store NEXT-Series-Number, into REGISTER R1
MOV @R1,#01h ;put DATA 01hex into MEMORY location pointed to by the address in REGISTER R1
MOV R2,#0Bh ;COUNTER to count down the 11 numbers in the Fibonacci Series, so put 0Bhex into R2
MOV A,#01h ;Using Accumulator to create number in series, initialize second number 01hex
BACK: ADD A,@R0 ;Create next number in series by adding Accumulator contents to last number created
      INC R0 ;Increment Pointer in R0 to point at number just created in series, in memory
      INC R1 ;increment Pointer in R1 to point where to put next number created in series
      MOV @R1,A ;... and store there, the next number in series just created in the Accumulator
      DJNZ R2,BACK ;Decrement COUNTER and loop back to LABEL "BACK" until count is down to Zero
HERE: SJMP HERE ;Just hang here

```

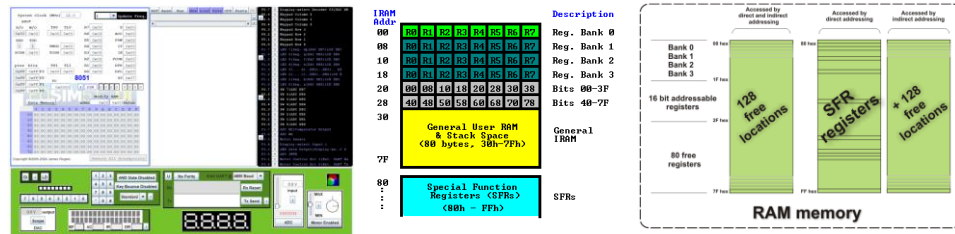
```

RTN: R0 <- 30h
RTN: (R0) <- 00h
RTN: R1 <- 31h
RTN: (R1) <- 01h
RTN: R2 <- 0Bh
RTN: A <- 01h
RTN: A <- A+(R0)
RTN: R0 <- R0+1
RTN: R1 <- R1+1
RTN: (R1) <- A
RTN: R2 <- R2-1, PC <- "BACK" address
RTN: PC <- "HERE" address

```

	0	1	2	3	4	5	6	7	8	9	A	B	C	D	E	F
00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00
10	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00
20	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00
30	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00
40	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00
50	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00
60	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00
70	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00	00

HINT (not given in Exam):



Recall this from both lecture and lab, and fill in the Register Transfer Notation (RTN):

Wunderlich, J.T. (1999). **Focusing on the blurry distinction between microprocessors and microcontrollers.** In *Proceedings of 1999 ASEE Annual Conference & Exposition, Charlotte, NC: (session 3547), [CD-ROM].* ASEE Publications. [PDF DOC](#)

Example 8051 microcontroller program using 8-bit arithmetic to do a 16-bit task; Decrement the 8-bit general-purpose registers R1 and R0 as one concatenated 16-bit number until it reaches the 16-bit number made by concatenating the contents of the 8-bit general-purpose registers R3 and R2.

```

check: MOV A, R0 ;put low-order byte in accumulator
      CJNE A, 02h, dcrmnt ;conditional jump to "dcrmnt" if not equal to R2 contents
      MOV A, R1 ;put high-order byte in accumulator
      CJNE A, 03h, dcrmnt ;conditional jump to "dcrmnt" if not equal to R3 contents
      SJMP done ;countdown finished, jump to "done"
dcrmnt: MOV A, R0 ;put low-order byte in accumulator
      CLR C ;must clear carry flag since used in subtraction
      SUBB A, #01h ;decrement (and possibly set borrow)
      MOV R0, A ;temporarily store new high-order byte in R0
      MOV A, R1 ;put high-order byte in accumulator
      SUBB A, #00h ;subtract borrow (i.e., carry bit is set if borrow at line #07)
      MOV R1, A ;temporarily store new high-order byte in R1
      SJMP check ;jump to "check"
done: NOP ;program finished

```

Example 8051 microcontroller program using 8-bit arithmetic to do a 16-bit task; Decrement the 8-bit contents of internal RAM addresses 21h and 20h as one concatenated 16-bit number until it reaches the 16-bit number made by concatenating the contents of the 8-bit general-purpose registers R3 and R2.

```

check: MOV A, 20h ;get low-order byte from on-chip RAM
      CJNE A, 02h, dcrmnt ;conditional jump to "dcrmnt" if not equal to R2 contents
      MOV A, 21h ;get high-order byte from on-chip RAM
      CJNE A, 03h, dcrmnt ;conditional jump to "dcrmnt" if not equal to R3 contents
      SJMP done ;countdown finished, jump to "done"
dcrmnt: MOV A, 20h ;get low-order byte from on-chip RAM for decrementing
      CLR C ;must clear carry flag since it is used as a borrow
      SUBB A, #01h ;decrement (and possibly set borrow)
      MOV 20h, A ;store new high-order byte in on-chip RAM
      MOV A, 21h ;get high-order byte from on-chip RAM for decrementing
      SUBB A, #00h ;subtract borrow (i.e., carry bit is set if borrow at line #07)
      MOV 21h, A ;store new low-order byte in on-chip RAM
      SJMP check ;jump to "check"
Done: NOP ;program finished

```

Wunderlich-IBM RESEARCH

List the FIVE main things in IBM S/390 & Z-Series PARALLEL PROCESSING (supercomputers)

worked on by J. Wunderlich Ph.D. Advisory Level Engineer & Researcher —MP4 YouTube

1. IMPLIMENTING 125 new IEEE BFP instructions (IBM previously used Hex Floating Point)
2. BRANCH PREDICTION with Branch History Table (BHT) —to cache branch addresses PDF
3. VIRTUAL ADDRESS PREDICTION with Translation Lookaside Buffer (TLB) —to cache Virtual Address Translations PDF
4. CACHE COHERENCY
5. QUALITY CONTROL / VERIFICATION with Wunderlich "Controlled Randomness" patented by IBM —PDF PPT

Wunderlich, J.T. (2003). Functional verification of SMP, MPP, and vector-register supercomputers through controlled randomness. In Proceedings of IEEE SoutheastCon, Ocho Rios, Jamaica, M. Curtis (Ed.): (pp. 117-122). IEEE Press. —PPTX w/audio PDF MP4 YouTube:

List the five main parts of this:

1. What does IID (Independent and Identically Distributed) mean in reference to Random Number Generators
2. What is a three-tuple graph in assessing the quality of a Random Number Generator (describe a bad one)
3. Discuss how Dr W's research of combining seven random number generators, and his developing API's for the IBM supercomputer test-kernel, provided varying options for the systems-level programmers at IBM (be specific)
4. List the three environments that the IBM supercomputer test-kernel ran in
5. Why did the selected Random Number Generators need to be run in reverse

WUNDERLICH IBM RESEARCH & DEVELOPMENT

—SOURCE: Dr. Wunderlich Lectures, AND —PDF MP4 YouTube

a. Why would IBM sell its supercomputer processors to Hitachi who were competitors? —ANSWER: to essentially "OWN" THEM if times got tough during competition

b. How did DUMPING of DRAM's onto the marketplace by South Korea during the early 1990s hurt Apple & Motorola, via IBM R&D?

ANSWER: IBM DRAM manufacturing and sales were funding the development of the IBM PowerPC chip (in Burlington Vermont) to go into Apple PC's TO KEEP APPLE AND MOTOROLA ALIVE SO THAT THEY COULD COMPETE WITH THE WINTEL MACHINES that were ironically never more than 10% of IBM revenue even during the 80's when IBM essentially owned the entire PC market (per Dr W two-hour private meeting with Dr Richard Attardi, general manager of IBM microelectronics and supervisor of 70,000 employees).

Why did IBM have an INTRANET? ... how was it constructed? ... and how did Dr. Wunderlich use it? —ANSWER: IBM in the early 1990s had its own INTRANET with OPTICAL FIBER LINES LEASED AND RUN BETWEEN ALL OF THE MAJOR CITIES IN THE UNITED STATES AND UNDER THE OCEAN TO EUROPE to PROTECT INTELLECTUAL PROPERTY; and how Wunderlich (remotely from new York) supervised a junior engineer in Austin TX who was implementing the quality control API's that Wunderlich developed (and that were patented by IBM since they owned his IBM intellectual property developments at that time), to run

WUNDERLICH PRE-IBM COMPUTER RELATED PROFESSIONAL EXPERIENCES

Discuss one non-IBM-related significant computing-related thing from Dr W's stories

ANSWERS:

1. Wunderlich's first neural network chip that he designed for his master's thesis (the picture on the wall of the lab is his second one) implemented mixed precision mathematics; that are now at the heart of deep learning neural network chips like those by NVIDIA that dominate the marketplace. Wunderlich did his own patent search at a microfiche patent depository library, then filed the patent disclosure document in the US Patent Office; however at that time it cost \$15,000 (in early 1990s dollars), for the attorneys to do the full-blown patent application, so the Wunderlich patent disclosure expired.
2. In the 1980's, when Wunderlich was in charge of building \$100 million (in 1980's dollars) worth of high tech office parks in Austin Texas and San Diego California, including DATA CENTERS, and needed to pull all of the 60 different contractors off of the 13 building office park in Texas, to fast track build a computer company's building in six weeks including a raised floor computer room that required extreme amounts of extra work to bring in the power for the cooling of the high performance water cooled computer systems.
3. Later, in La Jolla California in the 1980's, he built an office building and adjacent factory for a computer company by negotiating a multi \$1,000,000 contract to get it done fast track with very specific contract terms that clarified everything that needed to be done, partially based on his experience in Austin Texas.
- c. on workstation class machines, running IBM's version of UNIX (called AIX), that were then used decades later for Watson, the first high performance computer developed for deep learning

SCALABILITY IN PARALLELL PROCESSING

Summarize with one paragraph, and sketches as needed, at least three scalability aspects of computing, discussed in this course

DYNAMIC SCHEDULING: OUT-OF-ORDER EXECUTION IN PARALLEL PROCESSING (XX points)

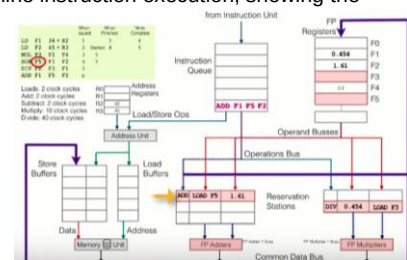
—SOURCES: PPTX PDF MP4 YouTube AND PPTX w/audio PDF MP4 YouTube

Make a sketch demonstrating TOMASULO'S ALGORITHM to facilitate out of order machine instruction execution, showing the RESERVATION STATION.

HINT:

Example: Dynamic scheduling using TOMASULO'S ALGORITHM to facilitate out of order machine instruction execution such that Add operation in the reservation station on the left will be executed before the Divide instruction in the reservation station on the right even though the divide operation is first in the instruction sequence. This is because the Divide operation requires 40 clock cycles to execute whereas the Add instruction only requires 2 clock cycles.

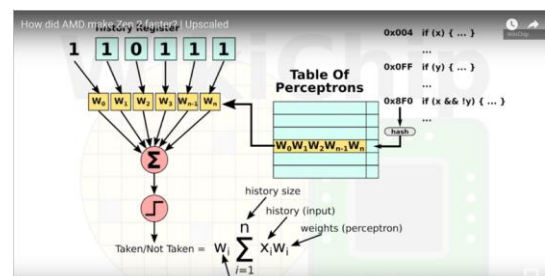
Watch: <https://www.youtube.com/watch?v=yjiE6NhtkiA>



(XX points) DYNAMIC SCHEDULING: OUT-OF-ORDER EXECUTION, USING NEURAL NETWORKS FOR PREDICTION

Make a sketch demonstrating how AMD makes Dynamic scheduling out of order execution faster using neural networks for prediction

HINT(not Given on Exam):



SOURCES:

- Describe how a **PULL-UP RESISTOR** works?
- Why can't you just splice in toggle switch or push button switch into an input line of any kind of device to presumably create logic zero and logic one?
- What is a **FLOATING-PIN** (i.e., why does every pin on an integrated circuit chip need to be connected to some voltage level if it is part of a digital circuit?)

SOURCES:

Why is it a good idea to use **RELAYS** to control outputs instead of trying to source the voltage and current directly out of the base unit?
Why use **TRI-STATE BUFFERS** to connect computer components to a bus?

SOURCES:

- CPU Design & Assembly Language intro [PDF](#)
- Computer Graphics ([PDF](#), [PPTX](#), [w/audio](#), [MP4](#), [YouTube](#))
- PC Design1 ([PDF](#), [PPTX](#), [w/audio](#), [MP4](#), [YouTube](#)), PC Design2 ([PDF](#), [PPTX](#), [w/audio](#), [MP4](#), [YouTube](#))
- Displays ([PDF](#), [PPTX](#), [w/audio](#), [MP4](#), [YouTube](#)), Graphics Boards ([PDF](#), [PPTX](#), [w/audio](#), [MP4](#), [YouTube](#))
- 2021 Etown GPU Design by J. Freaney ([YouTube](#))
- 2020 Etown CPU Design by J. Freaney and E. Schneider a ([MP4](#), [YouTube](#)), b ([MP4](#), [YouTube](#))
- 2016 Etown Super-Scalar Dual-Core Processor ([YouTube](#))
- SC-16H CPU Design by Glen G. Langdon ([PPTX](#), [w/audio](#), [PDF](#), [MP4](#), [YouTube](#))

Name three things that are parallel in a CPU or on a motherboard, and then name two things, other than what you just wrote, that are parallel in a GPU.

- Why is the penalty for an unexcused absence in this type of class so harsh?
- Why are teams often better than individuals working solo even if the net productivity is less?
- If you were to handpick a team of six people to work together under the most extreme of circumstances, define your selection criteria including what your role would be.
- What are the advantages of working with random potentially-broken equipment with potentially no solution?
- What happened on Apollo 13 and how did they survive?
- How could exercising on a secret military base compromise security? <https://www.theguardian.com/world/2018/jan/28/fitness-tracking-app-gives-away-location-of-secret-us-army-bases>
- Imagine and define a peacetime sustainability application of any of the technologies you have learned in this course.
- Imagine and define a military application of any of the technologies you have learned in this course.
- Imagine a new technology that you might create based on what you have learned in this course.
- What are the advantages of learning HDL Hardware Descriptive Language for programming FPGA's?
- What would be the downside of only learning HDL Hardware Descriptive Language to design digital circuits?
- Write in HDL Hardware Descriptive Language routine to implement a given digital logic design problem.
- Describe the use of timing diagrams using an FPGA
- Compare and contrast why you would use and FPGA versus Logisim for a given application
- What is the difference between the final "Synthesized" (not "optimized!") schematic created by the 2019 FPGA and the RTL Register Transfer Level schematic that you can generate, and how does this possibly relate to schematics made with Logisim. Describe how FPGA's in general, are more about the power and flexibility of implementing sophisticated digital design than about low power circuit design, or optimized logic
- Describe precisely what you would need to do to make a Logisim circuit that you create work on the 2019 FPGA. Compare and contrast digital design using the FPGA versus Logisim
- Define precisely what an "FPGA" Cell and LookUp table (LUT) are
- What does it mean to have a floating input or pin, and why is it a bad thing?
- Describe how a pull-up resistor works?
- Describe how a current-limiting resistor works with an LED?
- Describe how to debounce a toggle switch using logic gates
- Describe how to debounce a pushbutton using an analog circuit
- Why can't you just splice in toggle switch or push button switch into an input line of any kind of device to presumably create logic zero and logic one?
- Why does every pin on an integrated circuit chip need to be connected to some voltage level if it is part of a logic circuit implementation?
- Why do integrated circuit chips need to have a power and ground, but your simulated circuits do not?
- How do you make a voltage divider?
- What is the difference between a single phase and a three phase AC motor? Describe using applications.
- What are the different types of AC motors?
- What are the different types of DC motors, and how do you control them?
- When can you use pulse width modulation to control a motor?
- What is an optical encoder, and when would you use one?
- Describe what a motor's stator, rotor, and commutator are?
- What is the purpose of brushes in some motors?
- How is a DC stepper motor similar to a three phase AC motor?
- Describe in words and draw an internal electrical schematic of how a DC stepper motor works.
- Describe an application for when a DC stepper motor would be the best choice.
- Why are AC motors often a better choice for industrial applications?
- Draw an H bridge, describe how it works, and state an alternative way to create directional control for a DC motor.
- What is the difference between open loop and close loop control?
- What is the difference between a stepper motor and a servo?
- Draw a closed loop circuit showing 120 V DC power source, 0.8 ohm transmission line, and a 100 amp load, then calculate the voltage delivered to the load and state why it may not be acceptable. Also calculate the input power created by the source.
- Calculate the transmission line Power loss if we were to generate 120 V DC and deliver it to a 100 amp load over a transmission line with a 0.8 ohm resistance.
- List the three advantages of AC power over DC power
- Draw a two line diagram showing 120 V AC generator and a 1:10 step up transformer
- Draw a single line diagram showing 120 V AC generator, a 1:10 step up transformer, a 1200 V AC transmission line, a 10:1 step down transformer, and a 100 amp load. Then calculate the voltage delivered to the load and make a statement about its appropriateness in magnitude.
- Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Describe what **POWER FACTOR (PF)** is in an AC circuit
- Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Describe what **REACTIVE POWER** is in an AC circuit
- Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES** (Describe how to use a bank of capacitors in a factory to correct a **POWER FACTOR**)
- Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, What is a Phasor in reference to AC circuits; give an example starting with an AC sine wave.
- Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Describe in words the difference between real power, reactive power, apparent power an AC circuit power distribution.
- Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Draw the power triangle showing Real, reactive, and apparent power, and then state what a lagging and leading power factor is and how it is effected by an inductive or capacitive load

52. Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, What is the difference between kVA and kVAR?
53. Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Why do you believe the power company penalizes customers financially for poor power factor?
54. Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Why do you believe factories typically have poor Power Factors?
55. Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Describe in words what three phase power is, and also draw the three phases with voltage sine waves.
56. Concerning **POWER for DATA CENTERS and SILICON FOUNDERIES**, Draw a picture and describe in words what a harmonic is when discussing AC power, or anything having to do with waves including music, mechanical vibration control, or architectural acoustics.
57. Why is understanding how our **POWER generation and distribution** works critical to understanding our vulnerability to **CYBER ATTACK**? And what degree of security do you believe PLCs provide in preventing cybersecurity attack if they were to be networked and connected over the Internet?
58. In the United States we have three separate **POWER GRIDS**, list three reasons that we connect them together, and then list one reason concerning **CYBER SECURITY** that they should stay separate.
59. What are the **HAZARDS** to the function of **electronics in an industrial setting**?
60. What voltage can you input or output into the Axioline PLC?
61. What is the advantage of having sinking outputs on the nanoLLC?
62. What is **SCADA**?
63. Why does the problem of system **STABILITY** and close loop control only apply to **PLCs and not SCADA**?
64. How do you connect relays to the inputs and outputs of an axioline PLC?
65. Compare and contrast simulation versus real time control, including the strengths and weaknesses of both, and why you may want them to be working concurrently together.
66. What exactly happened at three Mile Island in 1979 and how did the people combined with their control systems fail?
67. Why do you think assembly language is used for embedded applications?
68. What did your group add to the Axioline PLC manual
69. Describe every part of a given Phoenix Contact data sheet for a relay
70. Describe polls and throws in a DPDT switch
71. Draw both the United States and international symbol for DPDT
72. Draw and describe in words exactly how you would connect wires to both the primary and secondary sides of a relay
73. Describe how to debounce a pushbutton using an assembly language routine
74. Describe in words how backpropagation neural network learning works
75. Describe the mobile robot real time and simulation assignments discussed in the Wunderlich publication **simulation versus real time control with applications to**
76. **robotics and neural networks**
77. Describe the robotic arm real time and simulation example discussed in the Wunderlich publication **simulation versus real time control with applications to robotics and neural networks**
78. Describe the Neural Network real time and simulation example discussed in the Wunderlich publication **simulation versus real time control with applications to robotics and neural networks**
79. Compare and contrast a raspberry pi to an Arduino
80. Compare and contrast why you would use a PLC versus a standard computer for a given application
81. Compare and contrast the Intel 8051 microcontroller to the ARM microcontroller, including both hardware and assembly language aspects
82. Compare and contrast PLC ladder logic to flowchart driven software implementation.
83. Describe two distinctly different aspects of Dr. Wunderlich's real time vs simulation of a bottling plant.
84. Describe two distinctly different ways that a simulation can be used in reference to developing a real time system (not two different applications, but rather two different purposes for one application)
85. Precisely describe one situation where an airplane pilot desperately needs the help of a simulation.
86. Precisely describe one situation when an airplane pilot needs to easily be able to override a simulation to take control of real time actuation's of aerodynamic control surfaces.
87. Precisely describe one situation when a fighter pilot makes the best use of concurrent interactive simulation and real time code.
88. Precisely describe what happened at three Mile Island and how the mix of simulation, real time code, and human screwup, led to what almost was the biggest environmental disaster in United States history. And describe how Chernobyl was different. And describe how Fukushima was different.
89. Why do you believe that the advanced PLC, unlike the nanoLC, does not seem to obviously have a way to simulate before implementation?
90. How could a cyber attack on water supply controls, including PLCs, of California, specifically effect southern California the most?
91. How could a cyber attack on the supply chain in the United States, including PLC controls of automated packing and distribution, affect the health of United States citizens (worst case)?
92. How could a cyber attack on a nuclear power facility, including PLCs, cause a disaster?
93. How could a cyber attack on an oil and gas facility, including PLC's, create the most damage?
94. How could a cyber attack on the United States electrical grid, including PLCs, cause the most damage?
95. What exactly happened with "Stuxnet"?
96. What manufacture of PLCs did our guest lectures tell us they use at GEA Systems?
97. What technology did our guest lectures from GEA systems demonstrate for us?
98. What was the most mission critical application of PLCs that our guest lecturers from GEA systems told us about?
99. Describe what an assembler directive is?
100. Describe the exact purpose of these four Intel 8051 assembler directives: "org" "equ" "db" "dw"
101. Why does the NOP instruction exist in 8051 microcontroller assembly language
102. Describe the **"NEURAL NETWORK PREDICTION"** and **"SMART PREFETCH"** of the AMD ZEN CORE 2017
103. (x points) From lecture **"Technology History & Economics"** How is the **NASDAQ** different from the **DOW**
104. List all computer-relevant Powers of 10 and 2 [PDF](#) -- no video of lecture, but [video review in #7 below](#). (All but **precise** powers of two column)
105. Chip Manufacturing Process [PDF](#) -- no video of lecture, but [video review in #7 below](#) (All of first diagram)
106. Atoms and Transistors [PDF](#) -- no video of lecture, but [video review in #7 below](#) (Doping, and how each type transistor works)
107. Moore's Law [PDF](#) -- no recording of lecture, but [video review in #7 below](#) (All, but no calculations)
108. PAPER: Transistors Stop Shrinking [PDF](#) -- part of an assignment, so no lecture or video, just class discussion (Not on exam)
109. BOOK CHAPTER: Computer History [PDF](#) -- part of an assignment, so no lecture or video, just class discussion (Not on exam)
110. Tech History & Economics [PDF \(MP4* YouTube*\)](#) -- [videos includes review of #1 to #4 above \(in the beginning\)](#) (All that is highlighted, except company names)
111. Sketch the Conceptual Computer Architecture [PDF \(MP4* YouTube*\)](#) -- [videos include #9 below](#) (All)
112. Levels of Computing, Microcontroller vs Microprocessors, Robotics, IBM quality control [PDF PPT \(MP4* YouTube*\)](#) -- [videos include #8 above](#)
113. (Only slides 1,8,9,10,11,12,13,16,18,19)
114. Microcontroller vs Microprocessors [PDF \(MP4* YouTube*\)](#) (All highlights 4) Wunderlich, J.T. (1999). **Focusing on the blurry distinction between microprocessors and microcontrollers**. In *Proceedings of 1999 ASEE Annual Conference & Exposition*, Charlotte, NC: (session 3547), [CD-ROM]. ASEE Publications. [PAPER](#)
115. IBM/Wunderlich "Controlled Randomness" Quality Control [PDF PPTX w/audio \(MP4* YouTube*\)](#) -- [EGR433/430 Advanced Computer Engr/Parallel Processing](#)
116. Wunderlich, J.T. (2003). **Functional verification of SMP, MPP, and vector-register supercomputers through controlled randomness**. In *Proceedings of IEEE SoutheastCon, Ocho Rios, Jamaica*, M. Curtis (Ed.): (pp. 117-122). IEEE Press. [PAPER](#) (All)
117. Intro to Cache Design [PDF \(MP4* YouTube*\)](#) --
118. Number Representations [PDF PPTX w/audio MP4 YouTube](#) (All)
119. Fractional part of IEEE Floating Point [PDF PPTX w/audio MP4 YouTube](#) (All)

120. IEEE Floating Point example [PDF PPTX w/audio MP4 YouTube](#) (All)
121. Design a PC 1 [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include #17 below* (All slides except 7,9,10,11,12,13)
122. Design a PC 2 [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include #16 above* (All slides except 4 to 12, and 15 to 18)
123. BOOK CHAPTER: RISC vs CISC, HLL vs Assembly [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include #19 below* (—)
124. Physics and Technology of Waves [PDF PPTX w/audio \(MP4* YouTube*\)](#) (All but last four questions on slide 10)
125. Human vs. Machine Vision [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include excerpts from #22 below* (Beta movement, 4 ways our eyes different from camera capture)
126. “Natural & Man-Made Lighting” from *EGR353 Green Architectural Engineering* [PDF PPTX w/audio MP4 YouTube](#) (Not on exam)
127. Physics of Color, and Display Technologies [PDF PPTX w/audio \(MP4* YouTube*\)](#) (All about RGB, CRT, LCD, and plasma)
128. BOOK CHAPTER: Computer Graphics [PDF PPTX w/audio MP4 YouTube](#) (—)
129. Graphics Boards [PDF PPTX w/audio MP4 YouTube](#) (All slides except 5, and 11 to 27)
130. BOOK CHAPTER: Memory [PDF PPTX w/audio MP4 YouTube](#) (—)
131. BOOK CHAPTER: Storage [PDF PPTX w/audio \(MP4* YouTube*\)](#) (—)
132. BOOK CHAPTER: Processors [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include all of #29,30,31,32 below* (—)
133. AMD ZEN core [PDF PPTX w/audio MP4 YouTube](#) Video: “How did AMD make Zen 2 faster?” (Know “Neural-Net Prediction” and “Smart Prefetch”)
134. Amdahl’s Law for Parallel Processing [PDF PPTX w/audio MP4 YouTube](#) (All)
135. PAPER: Breaking Multicore Bottleneck [PDF PPTX w/audio MP4 YouTube](#) (Explain and graph how it relates to Amdahl’s Law)
136. Recent Intel microprocessors [Wikipedia](#) Video: “Intel’s new processors and GPUs in under 10 minutes | CES 2022” (—)

137. Routers [PDF PPTX w/audio MP4 YouTube](#) (Difference between router and modem)
138. Clean Power 1 [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include #35 below, plus notes on Power Factor* (—)
139. Clean Power 2 [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include #34 above, plus notes on Power Factor* (Page 1: what transformer and rectifier do. Page 7: six things about clean power)
140. Intro to Machine Intelligence Symbolic AI vs Neural Networks [PDF PPTX w/audio \(MP4* YouTube*\)](#) *—videos include #37 below, plus video on Machine-Learning Math* (—)
141. Neural Network Code runs (part of my 1991 Neurocomputer chip development) [MP4 YouTube](#) (Not on exam)
142. Virtual & Augmented Reality [PDF PPTX w/audio MP4 YouTube](#) (State one good, and one bad thing from each of these three papers)
143. 2020 Etown Oculus Rift VR of Campus in 1924 and Present, and in Revit, by JJWIV [YouTube](#) (Not on exam)
144. Human vs Machine Intelligence [PDF PPTX w/audio \(MP4* YouTube*\)](#) (Explain difference between Symbolic AI and Neural Networks)
145. Technology and Humanity III [PDF PPTX w/audio MP4original YouTube : with class discussion: \(MP4* YouTube*\)](#) (Defend one technology that needs minimum regulation, and one that needs maximum regulation/limits, and why) Also watch Star Trek “Measure of a Man” or “Brothers” (Possible extra credit question related to the android “Data”)
146. Robotics & Machine Intelligence at Elizabethtown College since 1999 [PDF PPTX MP4 YouTube](#) (Not on exam)
147. Guest Lecture by U.S. Ambassador John B. Craig on U.S. National Security Policy [YouTube](#) (—)
148. Elizabethtown College “Ware Lecture” Cybersecurity Symposium, with three of world’s top experts [YouTube](#) (—)
149. 101 Questions compiled from Students and Faculty [PDF](#) (Not on exam)
150. Where is the company our guest speakers are from?
151. Approximately how many billions of dollars of revenue and how many employees work at our guest speakers company?
152. Approximately how many of our computer engineers have been employed at this company?
153. Name one aspect of our guest speakers technology that could be very sensitive to cyber attack?
154. Name one application of our guest speakers technology that requires extremely high precision temperature control?
155. Name one application of our guest speakers technology that is most dangerous if not controlled carefully?
156. Write a paragraph about what our guest speakers talked about.
157. Describe how a design-build turnkey company works
158. What manufacture of PLCs does our guest speakers company use?
159. Describe at least one component of a thermodynamic circuit in a refrigerator.
160. Write a paragraph about what you would hope to see and learn on a field trip to our guest speakers’ work.

d.—