



סילבוס שם המחלקה:

Course Title	חומרה מאיצה	שם הקורס
.Course No	3624550-20	מספר קורס
Lecturers' Name	אורן גנון	מרצי הקורס
Year	2025-2026	שנת הוראה
Weekly Hours	2 שעות הרצאה + 2 שעות תרגול	היקף הקורס בש"ש
Credits	3	נקודות זכות
Prerequisite	מערכות סיפרתיות תכן של חומרת מחשבים	דרישות קדם

תקציר הקורס

מטרת הקורס היא לחשוף את הסטודנטים לעולם המאצים החישוביים המקביליים, ובפרט – GPU – Graphics Processing Unit כמאיץ מקבילי רב-עוצמה. במהלך הקורס הסטודנטים ירכשו ידע תיאורטי ומעשי על עקרונות הארכיטקטורה של GPUs, המודלים של עיבוד מקבילי והאינטראקציה שלהם עם מערכת ההפעלה והמעבד המרכזי (CPU).

הקורס שם דגש על Hands-On Programming: הסטודנטים יכתבו קוד ב-CUDA כדי ליישם בפועל את העקרונות הארכיטקטוניים שנלמדים בכיתה, יתנסו בניהול משאבים על GPU, יבצעו אופטימיזציות לביצועים, ולבסוף יפתחו מערכת מואצת שמשלבת תאוריה עם יישום חישובי אמיתי.

Course Summary

The goal of this course is to introduce students to the world of **parallel computing accelerators**, with a special focus on the **GPU – Graphics Processing Unit** as a powerful parallel accelerator. Students will gain both theoretical and practical knowledge about GPU architectures, parallel processing models, and their interaction with the operating system and the CPU.

The course emphasizes **Hands-On Programming**: students will write CUDA code to implement the architectural concepts learned in class, practice GPU resource

management, apply performance optimizations, and ultimately develop an accelerated system that combines theoretical foundations with real-world computational applications.

אופני הוראה עיקריים

שיעורי הרצאות פרונטליות יחד עם סדנה ותרגול מעבדה
הקורס הוא היברידי כלומר בחלקו מקוון בחלקו פנים-אל-פנים /

תוצאות הלמידה

הגדרת תוצאות למידה: מה שהסטודנטיות ים ידעו או יוכלו לעשות בסיום הקורס

הבנה תאורטית

- להסביר את העקרונות המרכזיים של עיבוד מקבילי והבדלים בין CPU ל-GPU.
- לתאר את מבנה הארכיטקטורה של GPU (כולל CUDA cores, SMs, היררכיית זיכרון).
- להבין את תפקיד מערכת ההפעלה ותשתיות התוכנה בתמיכה במאיצים מקביליים.

כישורים טכניים

- לכתוב קוד בסיסי ומתקדם ב-CUDA תוך שימוש נכון ב-Threads, Blocks ו-Grids.
- לנהל ולהעביר נתונים בין CPU ל-GPU ולבצע שימוש יעיל בזיכרון משותף וגלובלי.
- לבצע פרופיילינג ואופטימיזציה לקוד GPU כדי לשפר ביצועים (memory coalescing, occupancy optimization).

יישום מעשי - בתרגולים

- ליישם אלגוריתמים חישוביים נפוצים (כגון כפל מטריצות, Reduction, Prefix Sum) בסביבת CUDA.
- לפתח מערכת מואצת מבוססת GPU המשלבת ידע תאורטי עם מימוש פרקטי.
- להעריך ביצועים ולנתח trade-offs בין פתרון סדרתי לפתרון מקבילי.

תוצאות למידה – מיני-פרויקט

בסיום המיני-פרויקט הסטודנטיות והסטודנטים יוכלו:

1. יישום מעשי של עקרונות CUDA

- o לממש אלגוריתמים חישוביים אמיתיים על גבי GPU באמצעות CUDA.
- o לבצע חלוקה נכונה של Threads, Blocks ו-Grids בהתאם לבעיה.

2. ניהול משאבים וזיכרון

- o לנהל ביעילות העברת נתונים בין CPU ל-GPU.
- o להשתמש במבני זיכרון שונים (Shared, Global, Constant) כדי לשפר ביצועים.

3. ביצוע אופטימיזציות

- o לזהות צווארי בקבוק בביצועים באמצעות כלי פרופיילינג.
- o ליישם שיטות אופטימיזציה (כמו Memory Coalescing או שימוש ב-Shared Memory) לשיפור ריצה.

4. מיומנויות חקר וחשיבה ביקורתית

- o להעריך Trade-offs בין פתרון סדרתי לפתרון מקבילי.
- o להשוות בין מימוש נאיבי למימוש מואץ ולנתח שיקולים של דיוק מול ביצועים.

5. עבודת צוות והצגה

- o לעבוד בצוות קטן לפתרון בעיה חישובית מקבילית.
- o להציג את הפרויקט, התוצאות והמדדים (ביצועים, דיוק, צריכת זיכרון) בצורה מסודרת ושקופה

מבנה הערכה		
קטגוריה	% מהציון	הערות
פרויקט גמר	70%	
תרגילי להגשה	30%	
	__%	
	__%	
סה"כ	100%	

חובות הסטודנטים	על פי תקנון שנקר.
-----------------	-------------------

מבנה הקורס			
	מספר/תארי	נושא השיעור	פירוט
	1	מבוא למאצים חישוביים	נלמד על ההבדלים בין CPU ל-GPU, הרקע לצורך במקביליות, חוק אמדל והשפעתו על ביצועים. בתרגול: חישובים סדרתיים מול מקביליים, הדגמה בסיסית של האצת ביצועים.
	2	מאצים מוכוונים תפוקה (Throughput-oriented Accelerators)	את עקרונות התכנון של GPU כמאיץ מקבילי מב-תפוקה. בתרגול: כתיבת תוכנית CUDA ראשונה, הרצת Kernel פשוט.
	3	אופטימיזציות ב-GPU	אזה: שיטות לניצול מקסימלי של ליבות GPU, latency על ידי SIMT scheduling. תרגול: נודה עם Nsight / Profiler – ניתוח ביצועים בסיסי (self-study מודרך).
	4	מודלי זיכרון ב-GPU	

רצאה: היררכיית הזיכרון (גלובלי, משותף, קומי, Registers). תרגול: מימוש אלגוריתם CU (חלק 1) תוך שימוש נכון במבני הזיכרון.			
מודל התכנות ו-API של CUDA	5	אח: ארגון Thread-Block-Grid, שימוש ב-API של CUDA לניהול משאבים. תרגול: מימוש אלגוריתמים נוספים (חלק 2) – למשל Vector Reduction או Add.	
זיכרון ולקליות של נתונים	6	רצאה: חשיבות locality, caching ושימוש ב-Shared Memory. תרגול: שימוש ב-cuda::atomic לפתרון מצבים של גישה מתחרה לנתונים.	
תבניות חישוב מקבילי: Convolution ושיקולי נקודה צפה	7	רצאה: מימוש קונבולוציה על GPU, אתגרי Floating-Point, דיוק מול ביצועים. תרגול: מימוש ל מטריצות ב-CUDA תוך שימוש ב-Shared Memory.	
הרצאת העשרה: ארכיטקטורות מתקדמות	8	אח: היכרות עם מאיצים נוספים (TPU, DPU, AI Accelerator), מגמות עכשוויות בתחום. ל: העמקה חופשית בנושאי CUDA שנבחרים ע"י הסטודנטים.	
תבניות חישוב – Parallel Prefix Sum (Scan)	9	אח: אלגוריתם Prefix Sum במימוש מקבילי. גול: מימוש האלגוריתם ב-CUDA תוך ניצול זיכרון משותף.	
חישוב היסטוגרמה מקבילי	10	רצאה: בעיות scaling, contention, ושיטות יעול חישוב היסטוגרמות על GPU. תרגול: מימוש היסטוגרמה מקבילית עם CUDA.	

אה: אתגרים של חישוב על מטריצות דלילות, צוגים שונים (CSR, COO). תרגול: התחלת ני-פרויקט – בחירת בעיה חישובית ויישום ראשוני ב-CUDA.	Sparse Matrix Computation	11	
Threading הרצאה: סקירת חומרת ה- Streams נושאים מתקדמים (כמו GPU-ב- תרגול: המשך וסיום. Events) מיני-פרויקט והצגת תוצאות	Threading Hardware ומבנים מתקדמים	12	
הרצאה: שילוב עבודה היברידית בין CPU ל-GPU, מודלים של חלוקת עומס, השוואה בין ביצועים ותיאום משאבים. היכרות עם ספריות כמו OpenACC או Thrust.	GPU + CPU Heterogeneous Programming	13	
הרצאה: סקירת היישומים המרכזיים של GPU בלמידה עמוקה – אימון רשתות ניורונים, פעולות בסיסיות (Convolution, Backpropagation). תרגול: יישום Forward Pass קטן ברשת ניורונים ב-CUDA, או שימוש בספרייה קיימת (cuDNN) כדי להריץ רשת פשוטה.	Deep Learning on GPUs	14	

רשימה ביבליוגרפית	
"Programming Massively Parallel Processors: A Hands-on Approach"	רשות
University of Oxford: CUDA Programming on NVIDIA GPUs Taught by Mike Giles, Professor	רשות
UC Davis: EE171: Parallel Computer Architecture Taught by John Owens, Associate Professor	רשות

University of Sheffield: COM4521: Parallel Computing with GPUs Taught by Paul Richmon	רשות