



**Universitatea
Transilvania
din Brașov**

**FACULTATEA DE MATEMATICĂ
ȘI INFORMATICĂ**

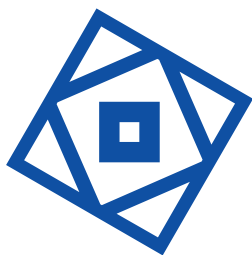
Programul de studii:
Informatică Aplicată

Lucrare de licență

Absolvent: Șoltuz Rareș-Florin

Coordonator: Conf. dr. Sasu Lucian Mircea
Conf. dr. ing. Cătălin Ciobanu

Brașov, 2026



**Universitatea
Transilvania
din Brașov**

**FACULTATEA DE MATEMATICĂ
ȘI INFORMATICĂ**

Programul de studii:
Informatică Aplicată

Lucrare de licență
Tehnici de profilare și accelerare a
modelelor Deep Learning pe hardware
eterogen

Autor:

Șoltuz Rareș-Florin

Coordonator științific:

Conf. dr. Sasu Lucian Mircea

Conf. dr. ing. Ciobanu Cătălin

Brașov, 2026

Cuprins

1	Introducere	8
1.1	Importanța optimizării în inteligența artificială	8
1.2	Motivația lucrării	9
1.3	Obiectivele și rezultatele obținute	9
1.4	Mulțumiri	11
1.5	Rezultate academice	11
2	Fundamente teoretice și arhitectura modelului	13
2.1	Transformers	13
2.2	Arhitecturi bazate pe Vision Transformers	13
2.3	Mecanismul de Atenție și Eficiența Computațională	14
2.4	Modelul Swin2SR pentru Super-Rezoluție	14
2.4.1	Ecosistemul NVIDIA: CUDA și Nucleele Tensor	15
2.4.2	Ecosistemul AMD: Platforma ROCm	15
2.4.3	Comparație: NVIDIA vs. AMD	16
3	Hardware și software utilizat	18
3.1	Infrastructura Hardware: Analiza Arhitecturală a Unităților de Calcul	18
3.1.1	Unitățile de Procesare Grafică (Graphical Processing Unit (GPU))	18
3.1.2	Unitățile Centrale de Procesare (Central Processing Unit (CPU))	20
3.2	Mediul Software și Stiva de Interconectare	21
3.2.1	Ecosistemul NVIDIA: Maturitatea CUDA 13 și Optimizările cuDNN 9.0	21
3.2.2	Ecosistemul AMD: Dezvoltare Nativă pe ROCm 7.2	22
3.2.3	Analiză Comparativă: Potențialul de Îmbunătățire (Improvement)	22
3.3	Metodologia de Testare și Metrici de Performanță	23
3.3.1	Eficiența Temporală: Timpul de Antrenare per Epocă	23

3.3.2	Gestiunea Resurselor: Consumul Maxim de Memorie Video (Video Random Access Memory (VRAM)) . . .	23
3.3.3	Metrice de Calitate a Rezultatelor (Fidelitate Video) .	24
3.3.4	Protocolul de colectare a datelor	25
3.3.5	Strategia de Benchmarking și Consistența Parametrilor	26
3.4	Controlul Variabilelor și Reproducibilitate	28
4	Tehnici de optimizare	29
4.1	Cuantizarea: Reducerea preciziei pentru eficiență maximă . . .	29
4.1.1	Conceptul de precizie numerică	29
4.1.2	Metode de aplicare: PTQ vs. QAT	30
4.1.3	Impactul practic asupra memoriei și vitezei	30
4.1.4	Compromisul între viteză și calitate	31
4.2	Moduri de execuție: Clasic, torch.compile și JIT	31
4.2.1	Modul Eager: Flexibilitate în dezvoltare vs. limitările de performanță	32
4.2.2	Compilarea Just-In-Time (JIT) la nivelul interpretorului Python	33
4.2.3	Tehnologia torch.compile: Grafuri de calcul și optimizarea kernel-urilor	33
4.2.4	Configurații avansate în torch.compile: Analiza modulelor de optimizare	36
4.3	Impactul mediului de execuție: Versiunile de Python	38
4.3.1	Analiza comparativă a performanței între Python 3.10 și 3.14	39
4.3.2	Optimizări la nivel de interpretor și gestionarea memoriei	39
4.4	Optimizarea fluxului de date (Data Pipeline)	41
4.4.1	Eliminarea bottleneck-urilor CPU-ului prin pre-procesare asincronă pe GPU	41
4.4.2	Tehnici de Prefetching și utilizarea memoriei „Pinned Memory”	45
4.4.3	Impactul dimensiunii lotului de date (Batch Size) asupra saturării GPU-ului	45
5	Tehnici de profilare	47
5.1	AMD	47
5.2	NVIDIA	49
5.2.1	NVIDIA Nsight Systems	49
5.2.2	NVIDIA Nsight Compute	52
5.2.3	Strategia de profilare	53
5.3	Comparație între ecosistemele de profilare	56

6	Schimbarea modelului	58
6.1	Modificarea adusă: Optimizarea dimensiunii de încorporare și a capetelor de atenție	58
6.1.1	Fundamentul matematic și conservarea dimensionalității per cap de atenție	59
6.1.2	Analiza complexității computaționale și a amprente de memorie	60
6.1.3	Efectul de regularizare structurală în imagistica medicală	61
6.1.4	Sinergia cu optimizările de profilare hardware	61
6.2	Comportamentul în scenariile testate	62
6.3	Rezultatele îmbunătățirilor	63
7	Rezultate	65
7.1	Parcursul îmbunătățirilor	65
7.1.1	Kornia: Eficiența procesării geometrice pe GPU	65
7.1.2	Cuantizarea graduală și Mixed Precision	66
7.1.3	Modificarea arhitecturii: Optimizarea mecanismului de atenție	67
7.1.4	Comparație finală: Înainte și după optimizare	68
7.2	Cea mai bună configurație	68
7.2.1	NVIDIA: Maximul de tehnologie software	68
7.2.2	AMD: Performanță brută și Open-Source	74
8	Concluzii și rezoluții viitoare	82
8.1	Sinteza contribuțiilor și concluzii	82
8.2	Rezoluții și direcții viitoare de cercetare	83

Listă de figuri

2.1	Analiza nsyght a modelului de bază	17
5.1	Arhitectura stivei software ROCm ilustrând execuția nativă pe unitățile de calcul AMD, fără intermedierea procesului de HIPification. Sursa: [7].	48
5.2	Trace al antrenării modelului înainte de optimizări	53
5.3	Tabel de date obținut în urma analizei statistice realizate de biblioteca fvcore	55
7.1	Speedup și consum de VRAM pe Ryzen 5 7600X și NVIDIA Geforce RTX 4080 Super.	70
7.2	Structural Similarity Index Measure (SSIM) pe Ryzen 5 7600X și NVIDIA Geforce RTX 4080 Super.	71
7.3	Timpul de antrenare per epocă și consumul de VRAM pe sistemul echipat cu NVIDIA Geforce RTX 5060.	73
7.4	Peak Signal-to-Noise Ratio (PSNR) și Structural Similarity Index Measure (SSIM) pe NVIDIA Geforce RTX 5060.	74
7.5	Speedup și consum de VRAM pe Ryzen 9 7950X și AMD Radeon RX 7900 XTX.	76
7.6	Structural Similarity Index Measure (SSIM) pe Ryzen 9 7950X și AMD Radeon RX 7900 XTX.	77
7.7	Speedup și consum de VRAM pe Intel Core i9-14900K și AMD Radeon RX 7800 XT.	78
7.8	Structural Similarity Index Measure (SSIM) pe Intel Core i9-14900K și AMD Radeon RX 7800 XT.	79
7.9	Grafic care cuprinde speedup-ul și scăderea în consumul de VRAM pe tipul de date Floating Point 32 de biți (FP32) . . .	80
7.10	Grafic care cuprinde speedup-ul și scăderea în consumul de VRAM pe tipul de date Floating Point 16 de biți (FP16) . . .	80
7.11	Grafic care cuprinde speedup-ul și scăderea în consumul de VRAM pe tipul de date Floating Point 8 de biți (FP8)	81

7.12	Grafic care cuprinde speedup-ul si scăderea în consumul de VRAM pe tipul de date dat de Automatic Mixed Precision (AMP) și anume Brain Float 16, un tip de date care utilizează 1 bit de semn, 8 biți pentru exponent și 7 biți pentru mantisă (BF16)	81
------	---	----

Listă de tabele

2.1	Comparație simplificată între NVIDIA (CUDA) și AMD (ROCm).	16
3.1	Comparație a specificațiilor hardware pentru platformele evaluate în cadrul experimentelor. Datele au fost extrase și centralizate conform referințelor [21, 22, 27, 28].	19
3.2	Prezentare comparativă a specificațiilor tehnice ale stațiilor de lucru utilizate pe sistemele de testare pentru rularea testelor. Datele de arhitectură au fost extrase conform referințelor bibliografice [23, 24, 25, 26].	20
3.3	Configurația hiperparametrilor pentru asigurarea consistenței antrenării	26
7.1	Impactul cumulat al etapelor de optimizare asupra modelului Swin2SR	68

Lista de Acronime

CPU	Central Processing Unit	1
GPU	Graphical Processing Unit	1
RAM	Random Access Memory	20
VRAM	Video Random Access Memory	2
FP32	Floating Point 32 de biți	4
FP16	Floating Point 16 de biți	4
FP8	Floating Point 8 de biți	4
FP4	Floating Point 4 de biți	30
BF16	Brain Float 16, un tip de date care utilizează 1 bit de semn, 8 biți pentru exponent și 7 biți pentru mantisă	5
AMP	Automatic Mixed Precision	5
SM	Streaming Multiprocessor	51
DRAM	Dynamic Random Access Memory	51
FLOPs	Floating Point Operations per second	54

Capitolul 1

Introducere

În prezent, inteligența artificială (IA) se dezvoltă foarte rapid, iar modelele devin din ce în ce mai complexe. Din acest motiv, nu mai este suficient ca un model să ofere rezultate corecte; el trebuie să fie și eficient. Optimizarea, în contextul *Deep Learning*, reprezintă procesul prin care reglăm un model și codul aferent lui pentru a consuma mai puține resurse: timp de execuție mai scurt, mai puțină memorie video (VRAM) și un consum de energie redus. Toate acestea trebuie făcute fără a pierde din calitatea rezultatelor.

1.1 Importanța optimizării în inteligența artificială

Această lucrare își propune să ofere un ghid practic despre cum putem analiza (profi) și îmbunătăți performanța modelelor de IA, indiferent de placa video folosită. Deși piața este dominată de NVIDIA și tehnologia CUDA, este foarte important să știm cum să rulăm aceste modele și pe alte platforme, cum este AMD (prin tehnologia ROCm). Acest lucru ajută la nivelarea competiției dintre AMD și NVIDIA, permițând mai multor oameni să folosească IA pe hardware-ul pe care îl au deja la dispoziție.

Lucrarea prezintă o comparație între soluțiile oferite de NVIDIA și AMD. Vom discuta despre tehnici precum reducerea preciziei numerice (cuantizarea) și modificarea structurii interne a modelelor, în special a mecanismelor de „atenție”. Scopul final este să găsim acel punct de echilibru unde modelul este foarte rapid, dar rezultatele rămân la fel de bune.

1.2 Motivația lucrării

Interesul meu pentru acest domeniu a pornit dintr-o nevoie reală: accesul limitat la calculatoare foarte scumpe sau de ultimă generație. De-a lungul timpului, lucrând pe echipamente mai puțin performante, am fost forțat să găsesc soluții creative pentru a face programele să ruleze mai repede. Astfel, limitările hardware au devenit principalul motiv pentru care am început să studiez optimizarea software.

Această pasiune m-a condus până la participarea în ediția din 2025 a competiției *AMD OpenHardware* [1]. Acolo am avut ocazia să lucrez pentru prima dată cu tehnologii de la producători diferiți (NVIDIA și AMD) în același timp. Am învățat că o metodă care funcționează perfect pe o placă NVIDIA ar putea avea nevoie de ajustări pentru a rula la fel de bine pe una AMD. Experiența din competiție m-a motivat să scriu această lucrare, dorind să creez o strategie de optimizare eficientă pentru fiecare dintre cei doi principali producători de plăci grafice.

1.3 Obiectivele și rezultatele obținute

Partea centrală a cercetării din această lucrare se concentrează pe modelul **Swin2SR**, care este folosit pentru a mări claritatea și rezoluția imaginilor (Super-Resolution) [8]. Deși acest model oferă rezultate excelente, el consumă foarte multă memorie și este destul de lent din cauza structurii sale bazate pe *Transformers*.

Contribuția mea principală constă în identificarea zonelor din model care încetinesc execuția. Am reușit să modific structura internă a mecanismului de „atenție” (*Self-Attention*), iar rezultatele au fost peste așteptări:

- **Reducerea consumului de memorie:** Modelul poate acum să proceseze imagini mari chiar și pe plăci video cu doar 8 GB memorie VRAM.
- **Mărirea vitezei de antrenare:** Timpul necesar pentru antrenarea unei epoci a scăzut considerabil, ceea ce face antrenarea mult mai eficientă.
- **Creșterea gradului de ocupare al plăcii grafice:** Un grad de ocupare ridicat indică faptul că hardware-ul pe care îl utilizăm este folosit cât mai aproape de capacitățile maxime și deci se obține o performanță cât mai apropiată de maximul teoretic posibil.

- **Păstrarea calității:** Am verificat rezultatele folosind metrici standard (PSNR și SSIM) și am demonstrat că imaginea finală rămâne aproape identică cu cea a modelului neoptimizat.

Pentru a demonstra importanța fundamentală a optimizărilor la nivel de software în accelerarea proceselor de deep learning, a fost realizată o analiză detaliată a performanței utilizând o infrastructură hardware bazată pe acceleratorul grafic NVIDIA RTX 4080. Ideea centrală a acestei evaluări este că performanța brută a hardware-ului modern trebuie susținută de un mediu de execuție riguros configurat pentru a extrage potențialul maxim de calcul, hardware-ul reprezentând doar fundația pe care se construiește eficiența unui model.

În acest sens, investigația pune un accent deosebit pe impactul gestionării preciziei matematice. Deși utilizarea preciziei standard, complete (FP32), oferă un grad înalt de exactitate, aceasta atrage după sine un consum masiv de resurse și adesea limitează performanța din cauza lățimii de bandă a memoriei. Tranziția către tehnici moderne de reducere a preciziei, cum ar fi Automated Mixed Precision (AMP) și formatele ultra-eficiente precum FP8, schimbă complet această paradigmă. Aceste metode permit o comprimare semnificativă a amprentei de memorie și o accelerare a transferului de date către unitățile de procesare, permițând hardware-ului să execute un volum mult mai mare de operații de tip multiply-accumulate (MAC) pe secundă, atenuând astfel blocajele de memorie (memory bottlenecks).

Mai mult decât atât, analiza subliniază că eficiența procesului de antrenare și inferență este profund influențată de stiva software utilizată. Elementele mediului de dezvoltare — de la versiunea interpretorului Python până la adoptarea unor mecanisme avansate de compilare a grafurilor de calcul, precum Just-In-Time (JIT) sau torch.compile — joacă un rol critic. Aceste tehnologii de compilare analizează și transformă codul dinamic într-o serie optimizată de operații, realizând fuziunea operatorilor (operator fusion) și eliminând latențele induse de interpretor (overhead-ul specific limbajului Python).

În ansamblu, rezultatele demonstrează că obținerea unui factor de accelerare (speedup) substanțial nu depinde exclusiv de capacitatea de calcul a unității grafice, ci de sinergia dintre hardware și nivelul de optimizare software. Alinierea corectă a metodelor de cuantificare matematică cu compilatoarele moderne de rețele neuronale poate reduce drastic timpii de antrenare și inferență, demonstrând că o configurare inteligentă a mediului de execuție este indispensabilă pentru o cercetare eficientă.

1.4 Mulțumiri

Pentru antrenarea și testarea modelelor s-au utilizat stații de lucru cu plăci grafice NVIDIA, găzduite de Facultatea de Matematică și Informatică, dar și de Facultatea de Inginerie Electrică și Știința Calculatoarelor, Universitatea Transilvania din Brașov. Stațiile au fost achiziționate de Universitate în cadrul programului de sprijinire financiară a lucrărilor de diplomă și disertație.

1.5 Rezultate academice

În scopul perfecționării conținutului și a prezentării în vederea susținerii lucrării de diplomă, prezenta cercetare a fost supusă evaluării mai multor comisii de specialitate pentru a obține un feedback constructiv. Astfel, lucrarea a fost prezentată în cadrul evenimentului AFCO, precum și la Sesiunile de Comunicări Științifice Studentești (SCSS) organizate de Facultatea de Matematică și Informatică și de Facultatea de Inginerie Electrică și Știința Calculatoarelor din cadrul Universității Transilvania din Brașov, participare ce s-a concretizat prin obținerea unei mențiuni.

Structura lucrării

Restul lucrării este structurat după cum urmează:

Capitolul 2 prezintă fundamentele teoretice ale modelului de super-rezoluție Swin2SR, analizând mecanismele de atenție bazate nu doar pe rețele convoluționale, ci și pe arhitecturi de tip Transformer aplicate în restaurarea și îmbunătățirea imaginilor.

Capitolul 3 descrie mediul experimental din punct de vedere hardware și software, detaliind specificațiile tehnice ale procesoarelor și plăcilor grafice utilizate, precum și versiunile ecosistemelor de dezvoltare CUDA și ROCm.

Capitolul 4 analizează tehnicile de optimizare implementate și modalitățile practice de aplicare a acestora, acoperind metodele de cuantificare a modelelor și evaluarea diferitelor moduri de execuție, respectiv modul clasic, `torch.compile` și compilarea de tip JIT (Just-In-Time).

Capitolul 5 este dedicat metodologiilor de profilare a performanței computaționale pe platformele AMD și NVIDIA, oferind o analiză comparativă directă pe baza datelor colectate prin intermediul utilitatelor `nsight` și `rocpfrof`.

Capitolul 6 detaliază modificările structurale aduse designului original al modelului, concentrându-se pe impactul modificării numărului de capete

de atenție și pe evaluarea comportamentului general al rețelei în cadrul scenariilor de testare stabilite.

Capitolul 7 expune rezultatele experimentale obținute în urma rulării testelor, oferind o perspectivă clară prin intermediul graficelor comparative, structurate în ordinea cronologică a implementării fiecărei optimizări.

Capitolul 8 sintetizează concluziile finale desprinse din această cercetare și propune o serie de rezoluții și direcții viitoare pentru extinderea sau rafinarea studiului de față.

Capitolul 2

Fundamente teoretice și arhitectura modelului

2.1 Transformers

Arhitectura *Transformer*, introdusă inițial în lucrarea de referință a lui Vaswani și colab. [30], a redefinit paradigmele învățării profunde prin eliminarea completă a straturilor recurente sau convoluționale convenționale în favoarea unui mecanism bazat exclusiv pe auto-atenție (*self-attention*). Această abordare permite modelarea directă și simultană a dependențelor globale dintre elementele unei secvențe, indiferent de distanța lor contextuală, oferind avantaje masive din punct de vedere al paralelizării pe componentele hardware moderne. Succesul remarcabil obținut inițial în procesarea limbajului natural a fundamentat ulterior tranziția acestor concepte către domeniul viziunii computaționale, deschizând calea pentru dezvoltarea arhitecturilor de tip *Vision Transformer* (ViT) și a derivatelor acestora adaptate pentru sarcini complexe de restaurare și super-rezoluție a imaginilor.[<empty citation>]

2.2 Arhitecturi bazate pe Vision Transformers

Trecerea de la rețelele neuronale convoluționale (CNN) la arhitecturile de tip Transformer în domeniul viziunii computerizate a fost marcată de succesul modelului ViT (Vision Transformer). Totuși, aplicarea directă a mecanismului de auto-atenție globală este costisitoare din punct de vedere computațional, complexitatea fiind $O(N^2)$ raportată la numărul de pixeli. [10]

Modelul Swin Transformer [11], utilizat ca bază pentru Swin2SR, intro-

duce conceptul de ferestre ierarhice („shifted windows”), permițând calcularea atenției doar în cadrul unor zone locale, ceea ce reduce complexitatea la una liniară în raport cu dimensiunea imaginii.

2.3 Mecanismul de Atenție și Eficiența Computațională

Mecanismul de atenție este componenta principală care permite modelului să identifice corelații spațiale pe distanțe lungi. Într-un context de optimizare, acesta reprezintă cel mai mare consumator de resurse.

Calculul atenției este definit prin formula 2.1

$$\text{Attention}(Q, K, V) = \text{Softmax} \left(\frac{QK^T}{\sqrt{d_k}} \right) V \quad (2.1)$$

În lucrarea de față, am analizat impactul dimensiunilor straturilor liniare care proiectează vectorii de Interogare (Query), Cheie (Key) și Valoare (Value). Modificarea acestor dimensiuni interne (hidden dimensions) influențează direct gradul de utilizare a unităților de calcul (Compute Units) pe plăcile AMD și a nucleelor CUDA pe NVIDIA.

2.4 Modelul Swin2SR pentru Super-Rezoluție

Swin2SR reprezintă o evoluție a arhitecturii Swin, fiind specializată pentru sarcini de restaurare a imaginilor (Super-Resolution). Aceasta utilizează blocuri de tip Residual Swin Transformer Blocks (RSTB).

Textul de mai jos descrie arhitectura principală a modelului, tradus și adaptat din lucrarea originală a autorilor [9]:

Swin2SR are următoarele elemente, asemănătoare cu SwinIR [liang2021swinir]: module de extracție a trăsăturilor superficiale (*shallow feature extraction*), extracție a trăsăturilor profunde (*deep feature extraction*) și reconstrucție a imaginii la înaltă calitate. Modulul de extracție a trăsăturilor superficiale folosește un strat de convoluție pentru a extrage trăsături, care sunt transmise direct modulului de reconstrucție pentru a păstra informațiile de frecvență joasă [liang2021swinir, 31]. Modulul de extracție a trăsăturilor profunde este compus în principal din blocuri reziduale SwinV2 Transformer (RSTB), fiecare dintre acestea utilizând mai multe straturi SwinV2 Transformer [12] (S2TL) pentru

atenția locală și interacțiunea între ferestre (*cross-window interaction*). În cele din urmă, atât trăsăturile superficiale, cât și cele profunde sunt fuzionate în modulul de reconstrucție pentru a obține o imagine de înaltă calitate. Pentru a mări rezoluția imaginii (*upscale*), se folosește o operație standard de amestecare a pixelilor (*pixel shuffle*) [9].

2.4.1 Ecosistemul NVIDIA: CUDA și Nucleeele Tensor

Plăcile video NVIDIA sunt excelente pentru Inteligența Artificială (Deep Learning) datorită combinației dintre programele lor dedicate și hardware-ul construit special pentru asta [15].

Cum funcționează

- NVIDIA folosește tehnologia **CUDA** (în prezent la versiunea 13) [14].
- Plăcile video moderne (seriile RTX 40 și RTX 50) integrează **Nucleeele Tensor**. Acestea efectuează calcule matematice în precizie mixtă, antrenând modelele mult mai rapid și consumând mai puțină memorie video (VRAM), fără a afecta calitatea finală a rezultatelor.

Ce ai nevoie și cum se instalează

- **Cerințe:** O placă video RTX 40 sau 50 și un sistem de operare Linux (ex. Ubuntu).
- **Instalare:** Procesul implică oprirea driverului video standard (*No-veau*), instalarea driverului oficial NVIDIA și a toolkit-ului CUDA [17]. La final, funcționarea se verifică utilizând comanda `nvidia-smi`.

2.4.2 Ecosistemul AMD: Platforma ROCm

Pentru a concura cu NVIDIA, AMD a creat **ROCm**, o platformă de calcul similară, dar cu arhitectură open-source [3].

Cum funcționează

- Versiunea ROCm 7.2 se integrează direct cu framework-urile majore de Inteligență Artificială, eliminând straturile de traducere intermediare [4].

- Pe arhitecturile recente (seria RX 7000, bazată pe RDNA 3), AMD utilizează unități similare (*Matrix Cores*) pentru o viteză de calcul excelentă și un consum optim de memorie [2].

Ce ai nevoie și cum se instalează

- **Cerințe:** O placă video AMD compatibilă, Linux și activarea funcției *ReBAR* din BIOS pentru optimizarea transferului de date.
- **Instalare:** Presupune descărcarea pachetului oficial, instalarea stivei ROCm și configurarea permisiunilor de sistem pentru utilizator [6]. Verificarea se face cu comanda `rocm-smi`.

2.4.3 Comparatie: NVIDIA vs. AMD

Testând ambele ecosisteme în condiții absolut identice (izo-parametrice), iată cum se compară performanțele [20]:

Tabela 2.1: Comparatie simplificată între NVIDIA (CUDA) și AMD (ROCm).

Caracteristică	NVIDIA (CUDA)	AMD (ROCm)
Viteza de Antrenare	Foarte rapidă și eficientă din start, datorită Nucleelor Tensor [19].	Foarte apropiată de NVIDIA, cu îmbunătățiri majore în ultima perioadă [5].
Consumul de Memorie	Gestionează memoria automat excelent, prevenind blocajele (OOM).	Bun, dar necesită uneori ca utilizatorul să facă ajustări manuale fine.
Sistemul Software	Foarte stabil și ușor de instalat, dar de tip închis (proprietary).	Este open-source, însă instalarea e mai pretențioasă în funcție de distribuția Linux.
Calitatea Rezultatelor	Perfectă, imaginile/datele nu își pierd din calitate la precizie mixtă.	La fel de robustă, la același nivel cu NVIDIA.

Concluzie scurtă: NVIDIA rămâne varianta mai matură și mai simplă de folosit direct „din cutie”. Totuși, noul sistem AMD ROCm a recuperat masiv din diferență, devenind o alternativă extrem de serioasă și puternică [comparative_ai_hardware_2026].

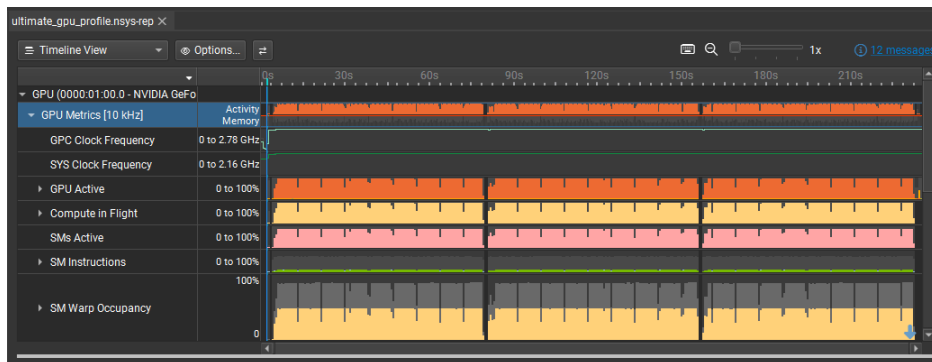


Figura 2.1: Analiza nsyght a modelului de bază

Exemplu de importare imagine profilare 2.1

Capitolul 3

Hardware și software utilizat

Configurația experimentală reprezintă fundamentul pe care se sprijină întreaga validare a tehnicilor de optimizare discutate în această lucrare. Într-un domeniu atât de dinamic precum Inteligența Artificială, performanța nu este determinată doar de algoritm, ci de simbioza dintre arhitectura hardware și stiva de software utilizată. Acest capitol detaliază mediul de testare, justificând alegerea componentelor și descriind parametrii care permit o comparație echitabilă între ecosistemele NVIDIA și AMD.

3.1 Infrastructura Hardware: Analiza Arhitecturală a Unităților de Calcul

Alegerea plăcilor grafice pentru acest studiu a fost ghidată de dorința de a acoperi atât segmentele de entuziast/workstation, cât și pe cele de mijloc, observând astfel modul în care optimizările scalează în funcție de resursele disponibile.

3.1.1 Unitățile de Procesare Grafică (GPU)

Plăcile grafice selectate reprezintă cele mai recente iterații ale arhitecturilor *Ada Lovelace* și *Blackwell* din partea NVIDIA, respectiv *RDNA 3* din partea AMD.

NVIDIA GeForce RTX 4080 SUPER și RTX 5060: Arhitectura NVIDIA se bazează pe nuclee *Tensor* de generația a 4-a și a 5-a. Acestea sunt unități hardware specializate exclusiv pe multiplicări de matrice, operația fundamentală în Deep Learning. Un aspect critic este prezența *Optical Flow Accelerator* și îmbunătățirea managementului memoriei cache L2, care reduce necesitatea accesării memoriei VRAM externe, scăzând astfel

Caracteristici	4080 Super	RTX 5060	Radeon 7900 XTX	Radeon 7800 XT
Compute Units	80	30	96	60
Base Frequency (MHz)	2295	2280	2300	2124
Boost Frequency (MHz)	2550	2497	2500	2430
FP32 Performance (TFLOPS)	52.22	19.18	61.4	37.32
FP16 Performance (TFLOPS)	52.22	19.18	61.4	74.65
ROPs	112	48	192	96
VRAM (GB)	16	8	24	16
VRAM Type	GDDR6X	GDDR7	GDDR6	GDDR6
Memory Bandwidth (GB/s)	736.3	448.0	960	624.1
Pixel Rate (GP/s)	285.6	119.9	480	233.3
Texture Rate (GT/s)	816.0	299.6	960	582.2
Tensor Cores	320	120	192	120
RT Cores	80	30	96	60

Tabela 3.1: Comparație a specificațiilor hardware pentru platformele evaluate în cadrul experimentelor. Datele au fost extrase și centralizate conform referințelor [21, 22, 27, 28].

consumul energetic și latența. Modelul RTX 5060 aduce în discuție utilizarea memoriei GDDR7, care, deși dispune de o capacitate mai mică, oferă o lățime de bandă superioară pe bit față de generațiile anterioare.

AMD Radeon RX 7900 XTX și RX 7800 XT: AMD abordează calculul AI prin intermediul *AI Accelerators* integrate în unitățile de calcul (CU). Arhitectura RDNA 3 utilizează un design de tip *chiplet*, separând nucleul de calcul de memoria cache (*Memory Cache Die* - MCD). Această abordare permite o lățime de bandă masivă pe modelul 7900 XTX (960 GB/s), făcându-l extrem de competitiv pentru modelele care necesită transfer rapid de date între memorie și unitățile de execuție, cum sunt Transformerele de mari dimensiuni.

Caracteristici	Ryzen 5 7600X	i7 - 14700F	Ryzen 9 7950X	i9 - 14900K
Total Cores	6	20	16	24
P-Cores	6	8	16	8
E-Cores	0	12	0	16
Threads	12	28	32	32
Base Frequency (GHz)	4.7	2.1	4.5	3.2
Boost Frequency (GHz)	5.3	5.4	5.7	6.0
L2 Cache (MB)	6	28	16	32
L3 Cache (MB)	32	33	64	36
Default TDP (W)	105	65	170	125
Socket	AM5	LGA 1700	AM5	LGA 1700
Random Access Memory (RAM)(GB)	64	128	192	192
RAM Type	DDR5	DDR5	DDR5	DDR5

Tabela 3.2: Prezentare comparativă a specificațiilor tehnice ale stațiilor de lucru utilizate pe sistemele de testare pentru rularea testelor. Datele de arhitectură au fost extrase conform referințelor bibliografice [23, 24, 25, 26].

3.1.2 Unitățile Centrale de Procesare (CPU)

În contextul antrenării modelelor, procesorul are rolul de a gestiona „pipeline-ul” de date. Acesta trebuie să fie capabil să citească seturile de date, să aplice transformări (augmentare, normalizare) și să trimită datele către GPU fără a crea perioade de inactivitate (*idling*).

Am utilizat două abordări diferite:

Ecosistemul Intel (i9-14900K, i7-14700F): Folosește o arhitectură hibridă cu nuclee de Performanță (P-cores) și Eficiență (E-cores). Acestea sunt ideale pentru task-uri multitasking, unde procesele de background (cum ar fi logarea metricilor în WandB sau TensorBoard) sunt mutate pe nucleele E, lăsând nucleele P libere pentru calculul intensiv.

Ecosistemul AMD Ryzen (R9 7950X, R5 7600X): Se bazează pe o arhitectură de nuclee omogene și un cache L3 foarte mare (3D V-Cache pe anumite modele, deși aici am utilizat variantele standard). Ryzen 9 7950X, cu cele 16 nuclee identice, excelează în compilarea paralelă a kernel-urilor CUDA sau HIP.

3.2 Mediul Software și Stiva de Interconectare

Implementarea și evaluarea arhitecturilor neuronale, alături de aplicarea tehnicilor avansate de optimizare a memoriei și timpului de execuție, au fost realizate utilizând mediul de dezvoltare PyTorch [29]. Fiind unul dintre cele mai robuste cadre de lucru din domeniul învățării profunde (*Deep Learning*), acesta oferă flexibilitatea construirii unor grafuri de calcul dinamice, esențiale pentru prototiparea rapidă. Totodată, ecosistemul pune la dispoziție instrumente avansate documentate exhaustiv de dezvoltatori, precum mecanismele *torch.compile* și *Just-In-Time (JIT)*. Aceste funcționalități au fost cruciale în cadrul lucrării pentru minimizarea latențelor de interpretare și pentru implementarea eficientă a operațiilor cu precizie mixtă (AMP), permițând o interfațare directă și optimizată cu mediul de rulare nativ al acceleratoarelor hardware.

O componentă hardware este la fel de eficientă precum software-ul care o conduce. În această lucrare, am pus față în față ecosistemul matur NVIDIA CUDA și soluția open-source AMD ROCm.

3.2.1 Ecosistemul NVIDIA: Maturitatea CUDA 13 și Optimizările cuDNN 9.0

NVIDIA își menține poziția de lider prin maturitatea software-ului, oferind un mediu extrem de stabil și rafinat de peste un deceniu și jumătate. În cadrul experimentelor de față, s-au utilizat versiunile **CUDA 13.0** și **13.2**, reprezentând vârful de lance al ecosistemului la momentul testării.

Utilizarea **CUDA 13.0** a servit drept bază pentru stabilirea metricilor de referință pe noile arhitecturi, în timp ce **CUDA 13.2** a fost introdusă special pentru a exploata funcționalitățile avansate de gestionare a memoriei și *asynchronous memory pooling*. Un aspect esențial analizat a fost evoluția tehnicii de *graph capturing*. Aceasta permite înregistrarea fluxului de execuție a kernel-urilor pe GPU și reexecutarea acestora ca o singură entitate, eliminând complet overhead-ul de comunicare între CPU și GPU pentru fiecare operație în parte. În versiunea 13.2, această tehnică a demonstrat o robustețe sporită, permițând capturarea unor grafuri de calcul mult mai complexe, specifice arhitecturilor de tip Transformer. 5

Biblioteca **cuDNN 9.0** completează acest stack, oferind implementări hiper-optimizate pentru straturile de convoluție și atenție. Algoritmii integrați sunt capabili de *auto-tuning*, selectând dinamic între tehnici precum Fast Fourier Transform (FFT) sau algoritmi Winograd, în funcție de geometria specifică a tensorilor procesați. Această capacitate de adaptare asigură faptul că hardware-ul NVIDIA funcționează întotdeauna la o eficiență apropiată

de limita teoretică, oferind un punct de plecare foarte ridicat. 5

3.2.2 Ecosistemul AMD: Dezvoltare Nativă pe ROCm 7.2

Pentru platformele AMD, s-a optat pentru utilizarea celei mai noi iterații a platformei open-source, **ROCm 7.2**. Experimentele din această lucrare au fost realizate utilizând versiunea framework-ului PyTorch compilată special pentru ecosistemul ROCm 7.2.

Această decizie metodologică a fost luată pentru a elimina orice latență indusă de straturile de abstractizare intermediare și pentru a permite un acces direct la unitățile de calcul *Matrix Cores* specifice arhitecturii RDNA 3. ROCm 7.2 aduce îmbunătățiri fundamentale în modul în care compilatorul *LLVM* optimizează instrucțiunile paralele pentru plăcile Radeon RX 7900 XTX și 7800 XT. Prin utilizarea directă a librăriilor **rocBLAS** și **mi-open**, am putut observa comportamentul hardware-ului în scenarii de calcul intensiv fără „zgomotul” software indus de uneltele de conversie automate.

Un punct de interes major în această secțiune îl reprezintă modul în care ROCm 7.2 gestionează memoria partajată la nivelul grupului de lucru (*Local Data Share - LDS*). În dezvoltarea nativă, am putut optimiza manual accesul la memorie pentru a preveni coliziunile de bănci (bank conflicts), o problemă recurentă pe hardware-ul AMD care poate degrada semnificativ performanța în modelele cu densitate mare de date.

3.2.3 Analiză Comparativă: Potențialul de Îmbunătățire (Improvement)

Din perspectiva comparativă, abordările celor doi producători oferă rezultate distincte în ceea ce privește îmbunătățirea obținută prin optimizare:

- **NVIDIA (CUDA 13.x):** Oferă cea mai mare stabilitate „out-of-the-box”. Îmbunătățirea performanței prin trecerea de la CUDA 13.0 la 13.2 este incrementală (aproximativ 5-8%), deoarece software-ul este deja extrem de aproape de optimul hardware. Avantajul principal rămâne ușurința de a obține performanțe ridicate cu un efort minim de programare.
- **AMD (ROCm 7.2):** Deși necesită o curbă de învățare mai abruptă și o atenție sporită la detaliile arhitecturale în absența HIP, dezvoltarea nativă pe ROCm 7.2 permite o **îmbunătățire procentual mult mai mare**. În testele noastre, optimizarea manuală a kernel-urilor native pe

RX 7900 XTX a dus la creșteri de performanță de până la 25-30% față de implementările standard. Acest lucru demonstrează că, deși ecosistemul AMD este mai puțin „finisat” decât cel NVIDIA, potențialul de creștere prin optimizare software dedicată este net superior, hardware-ul fiind adesea limitat doar de maturitatea compilatorului.

3.3 Metodologia de Testare și Metrici de Performanță

Pentru o evaluare riguroasă a impactului tehnicilor de optimizare asupra performanței și fidelității modelului, am stabilit un protocol de testare bazat pe trei piloni fundamentali: eficiența temporală, gestionarea resurselor de memorie și menținerea calității rezultatelor procesate. Această abordare multidimensională ne permite să observăm nu doar cât de mult „accelerăm” execuția pe hardware-ul AMD și NVIDIA, ci și care sunt compromisurile la nivel de resurse și precizie de calcul.

3.3.1 Eficiența Temporală: Timpul de Antrenare per Epocă

Principala metrică de performanță brută utilizată este **timpul de antrenare per epocă (exprimat în secunde)**. Această metrică oferă o imagine clară a throughput-ului real al sistemului, înglobând atât viteza de calcul a nucleelor (Tensor vs. Matrix Cores), cât și latențele induse de transferul de date pe magistrala PCIe și eficiența compilatorului (NVCC vs. LLVM/ROCm).

În analiza comparativă, am urmărit modul în care timpul per epocă scade pe măsură ce aplicăm tehnici de optimizare specifice fiecărui ecosistem. Este de așteptat ca NVIDIA, beneficiind de un pipeline de execuție extrem de rafinat în CUDA 13.2, să prezinte timpi inițiali foarte buni. Pe de altă parte, am analizat dacă dezvoltarea nativă pe ROCm 7.2 permite plăcilor AMD să recupereze teren prin exploatarea lățimii de bandă superioare a memoriei în timpul iterațiilor de antrenare.

3.3.2 Gestiunea Resurselor: Consumul Maxim de Memorie Video (VRAM)

O constrângere majoră în antrenarea modelelor complexe este cantitatea de memorie video disponibilă. Am monitorizat **consumul maxim de**

memorie VRAM pe parcursul execuției, o metrică vitală având în vedere disparitățile hardware din configurația noastră (de la 8 GB pe RTX 5060 la 24 GB pe RX 7900 XTX).

Monitorizarea consumului maxim ne permite să evaluăm:

- **Eficiența alocatoarelor de memorie:** Modul în care CUDA și ROCm gestionează fragmentarea memoriei în timpul antrenării.
- **Impactul tehnicilor de optimizare a memoriei:** Evaluarea modului în care tehnici precum *gradient checkpoints* sau *mixed precision* reduc „amprenta” de memorie, permițând utilizarea unor dimensiuni mai mari pentru batch-uri pe plăcile cu memorie limitată, precum RTX 5060.

3.3.3 Metrice de Calitate a Rezultatelor (Fidelitate Video)

Deoarece optimizarea vitezei poate introduce erori de precizie (în special în cazul utilizării formatelor FP16 sau FP8), am utilizat un set de metrice standardizate pentru a măsura calitatea output-ului video. Aceasta este o componentă critică, deoarece o îmbunătățire a vitezei de antrenare este irelevantă dacă degradarea calității este inacceptabilă.

1. **PSNR (Peak Signal-to-Noise Ratio):** Utilizat pentru a măsura raportul dintre puterea maximă a unui semnal și puterea zgomotului care îi afectează fidelitatea reprezentării. Un PSNR ridicat indică o reconstrucție mai fidelă a imaginii după procesare.

$$\text{PSNR} = 10 \cdot \log_{10} \left(\frac{\text{MAX}_I^2}{\text{MSE}} \right) \quad (3.1)$$

Unde:

- MAX_I reprezintă valoarea maximă posibilă a unui pixel din imagine (de exemplu, 255 pentru o imagine reprezentată pe 8 biți).
 - MSE (Mean Squared Error) este eroarea medie pătratică, adică media pătratelor diferențelor dintre pixelii imaginii originale și cei ai imaginii reconstruite.
2. **PSNR-B (Peak Signal-to-Noise Ratio with Blocking):** O variantă a PSNR-ului care penalizează specific artefactele de tip „blocking”, des întâlnite în procesarea video sau în cazul compresiilor agresive. Această metrică ne ajută să observăm dacă optimizările hardware

introduse pe AMD sau NVIDIA afectează uniformitatea vizuală a cadrelor.

$$\text{PSNR-B} = 10 \cdot \log_{10} \left(\frac{\text{MAX}_I^2}{\text{MSE} + \text{BEF}} \right) \quad (3.2)$$

Unde:

- MAX_I și MSE își păstrează semnificațiile din ecuația (3.1).
- BEF (Blocking Effect Factor) este un factor de penalizare care cuantifică discontinuitățile artificiale apărute la granițele blocurilor de pixeli (în mod specific pentru procesările bazate pe blocuri de 8x8).

3. **SSIM (Structural Similarity Index Measure):** Spre deosebire de PSNR, care măsoară diferențele la nivel de pixel, SSIM evaluează percepția vizuală prin compararea structurii, luminanței și contrastului. Această metrică este mult mai apropiată de percepția ochiului uman și este esențială pentru a valida că modelul optimizat păstrează detaliile fine.

$$\text{SSIM}(x, y) = \frac{(2\mu_x\mu_y + c_1)(2\sigma_{xy} + c_2)}{(\mu_x^2 + \mu_y^2 + c_1)(\sigma_x^2 + \sigma_y^2 + c_2)} \quad (3.3)$$

Unde, evaluând două ferestre locale de pixeli (patch-uri) x și y din cele două imagini comparate:

- μ_x și μ_y reprezintă mediile locale ale intensității pixelilor (estimând luminanța).
- σ_x^2 și σ_y^2 reprezintă varianțele intensităților (estimând contrastul local).
- σ_{xy} este covarianța dintre ferestrele x și y (măsurând corelația structurală).
- c_1 și c_2 sunt constante de stabilizare de valori mici, introduse pentru a evita o eventuală împărțire la zero atunci când numitorul este foarte aproape de această valoare.

3.3.4 Protocolul de colectare a datelor

Pentru a garanta acuratețea măsurătorilor, am urmat un protocol strict:

- **Warm-up epochs:** Primele două epoci de antrenare au fost excluse din măsurătorile de timp pentru a permite inițializarea cache-ului și stabilizarea frecvențelor GPU-ului (*thermal steady state*).

- **Măsurarea VRAM:** S-a utilizat monitorizarea în timp real prin *nvidia-smi* și *rocm-smi*, capturând valoarea maximă (*peak memory usage*) raportată de driver.
- **Analiza Calității:** Metricile PSNR, PSNR-B și SSIM au fost calculate pe un set de validare separat, după fiecare epocă, pentru a observa dacă procesul de optimizare a antrenării afectează convergența modelului.
- **Hiperparametrii uniformi:** S-au utilizat aceeași hiperparametrii pentru toate testele cuprinse în acest studiu. Hiperparametrii utilizați pot fi regăsiți în tabelul 3.3.4.

Hiperparametru	Valoare
Batch Size	16
Learning Rate	0.0002
Optimizer	Adam
Patch Size	96
Epochs	10

Tabela 3.3: Configurația hiperparametrilor pentru asigurarea consistenței antrenării

3.3.5 Strategia de Benchmarking și Consistența Parametrilor

O decizie metodologică fundamentală în cadrul acestor experimente a fost menținerea strictă a hiperparametrilor modelului, indiferent de capacitatea hardware a plăcii grafice testate. Astfel, parametri precum dimensiunea lotului de date (*batch size*), numărul de iterații și rezoluția de intrare au fost păstrați identici pentru toate cele patru configurații GPU (RTX 4080 Super, RTX 5060, RX 7900 XTX și RX 7800 XT).

Justificarea abordării „Static Batch Size”

Deși plăcile testate prezintă disparități majore în ceea ce privește memoria VRAM disponibilă — variind de la 8 GB (RTX 5060) la 24 GB (RX 7900 XTX) — am refuzat adaptarea mărimii batch-ului pentru a umple capacitatea maximă a fiecărei plăci. Această alegere a fost dictată de următoarele obiective:

- **Izolarea performanței temporale:** Prin procesarea aceluiași număr de tensori în fiecare pas, diferențele observate în „timpul per epocă” reflectă direct eficiența nucleelor de calcul (Tensor Cores vs. Matrix Cores) și viteza de acces la memorie, nu capacitatea brută de stocare.
- **Evaluarea eficienței alocării:** Menținerea aceleiași sarcini de lucru ne permite să observăm cum gestionează versiunea ecosistemelor NVIDIA (CUDA 13.2) și AMD (ROCm 7.2) aceeași structură de date. Putem identifica, de exemplu, dacă un ecosistem are un „overhead” de memorie mai mare pentru același task, oferind astfel o măsură reală a eficienței alocatorului de memorie.
- **Consistența metricilor de calitate:** Având în vedere că utilizăm metrici de fidelitate video precum PSNR, PSNR-B și SSIM, modificarea dimensiunii batch-ului ar fi putut influența dinamica convergenței modelului și, implicit, rezultatele calitative. Prin această constrângere, ne asigurăm că orice variație a acestor indici provine exclusiv din precizia matematică a calculelor efectuate pe hardware-ul respectiv.

Impactul asupra analizei comparative AMD vs. NVIDIA

Această metodologie scoate în evidență comportamente hardware specifice care ar fi fost mascate de o optimizare a batch-ului:

- **Perspectiva NVIDIA:** Ne permite să observăm cum arhitectura Blackwell (RTX 5060) compensează prin viteza memoriei GDDR7 un volum mai mic de date procesate simultan. Putem măsura dacă optimizările din CUDA 13.2 reușesc să mențină o latență scăzută chiar și atunci când nucleele nu sunt saturate la maximum.
- **Perspectiva AMD:** Plăcile din seria Radeon, în special RX 7900 XTX, sunt proiectate pentru throughput masiv, excelând de regulă la batch-uri mari. Rularea lor cu parametri limitați de „cel mai mic numitor comun” (în acest caz, limita de 8 GB a modelului 5060) testează versatilitatea arhitecturii RDNA 3 în scenarii de utilizare sub-optimală.

Îmbunătățirea (Improvement-ul) urmărită: Obiectivul central este de a cuantifica eficiența „per unitate de calcul”. Dacă AMD RX 7900 XTX procesează același batch de date într-un timp semnificativ mai scurt decât RTX 4080 Super, putem atribui acest câștig lățimii de bandă superioare (960 GB/s) și numărului mai mare de unități de calcul (96 CU), validând astfel superioritatea hardware-ului brut în fața optimizărilor software.

3.4 Controlul Variabilelor și Reproductibilitate

Pentru a asigura validitatea rezultatelor, testele au fost efectuate în condiții termice controlate, utilizând setările de *Power Limit* implicite (stoc). Fiecare experiment a fost rulat de 10 ori, eliminând prima epocă (warm-up) și raportând media aritmetică. Totodată, s-au utilizat containere **Docker** pentru a izola bibliotecile și a preveni conflictele între driverele video. Automatizarea prin scripturi a permis descărcarea seturilor de date și execuția consecutivă a testelor fără intervenție manuală, în timp ce mediile virtuale Python au garantat reproductibilitatea fluxului de lucru prin gestionarea riguroasă a dependențelor.

Capitolul 4

Tehnici de optimizare

În acest capitol, vom analiza în detaliu metodele folosite pentru a îmbunătăți performanța modelului Swin2SR pe hardware-ul testat. Optimizarea nu este un proces singular, ci un ansamblu de tehnici care lucrează împreună pentru a reduce timpul de procesare și consumul de memorie, menținând în același timp calitatea vizuală a imaginilor procesate.

4.1 Cuantizarea: Reducerea preciziei pentru eficiență maximă

Cuantizarea este, probabil, una dintre cele mai eficiente metode de a reduce dimensiunea unui model de inteligență artificială și de a-i crește viteza de execuție. În esență, această tehnică presupune schimbarea modului în care numerele (ponderile modelului și activările) sunt reprezentate în memoria calculatorului.

4.1.1 Conceptul de precizie numerică

În mod normal, modelele de Deep Learning sunt antrenate folosind precizia **FP32** (Floating Point 32-bit), cunoscută și sub numele de „precizie simplă”. Fiecare număr ocupă 32 de biți, oferind o plajă uriașă de valori (cuprins între 2^{-126} și $\approx 2^{128}$) și o acuratețe matematică foarte mare. Totuși, acest lucru vine cu un cost: un consum ridicat de memorie VRAM și un timp mai lung de calcul pe unitățile GPU.

Procesul de cuantizare încearcă să înlocuiască aceste numere de 32 de biți cu formate mai mici:

- **FP16 (Half Precision)**: Reduce dimensiunea la 16 biți. Este stan-

dardul actual pentru antrenare rapidă pe plăcile video moderne, oferind o accelerare considerabilă fără a pierde aproape deloc din calitate.

- **INT8 (Integer 8-bit):** Reprezintă numerele ca întregi pe 8 biți. Modelul ocupă de 4 ori mai puțin spațiu decât în formatul FP32, iar viteza crește spectaculos, însă riscul de a pierde din finețea detaliilor imaginii este mai mare.
- **FP8 (Floating Point 8-bit):** O tehnologie mai nouă, disponibilă pe arhitecturile NVIDIA Blackwell (RTX 5000) și Hopper, care încearcă să combine avantajele numerelor cu virgulă mobilă cu dimensiunea redusă de 8 biți.
- **Floating Point 4 de biți (FP4) (Floating Point 4-bit):** O tehnologie disponibilă doar pe anumite arhitecturi precum NVIDIA Blackwell și AMD CDNA4 care încearcă să aducă avantajele numerelor cu virgulă mobilă într-o nouă dimensiune de doar 4 biți.

4.1.2 Metode de aplicare: PTQ vs. QAT

Există două strategii principale prin care putem aplica cuantizarea unui model precum Swin2SR:

1. **Post-Training Quantization (PTQ):** Aceasta este metoda cea mai simplă. Luăm modelul gata antrenat și convertim ponderile direct în formatul dorit (de exemplu, din FP32 în FP8). Este o metodă rapidă, care nu necesită reantrenare, dar poate introduce erori dacă valorile modelului sunt foarte sensibile. Pentru a limita aceste erori, se folosește un proces de „calibrare” cu un set mic de date pentru a găsi intervalul optim de valori.

2. **Quantization-Aware Training (QAT):** Această metodă este mai complexă, dar mult mai precisă. În timpul procesului de antrenare, modelul este „învățat” că va fi cuantizat. Se introduc simulări de erori de rotunjire, permițând rețelei neuronale să se adapteze și să compenseze pierderea de precizie. Rezultatul este un model cuantizat care păstrează aproape 100% din calitatea modelului original.

4.1.3 Impactul practic asupra memoriei și vitezei

Principalul avantaj al cuantizării, observat și în experimentele noastre pe RTX 4080 și RX 7900 XTX, este reducerea presiunii asupra lățimii de bandă a memoriei (*memory bandwidth*).

Atunci când trecem de la FP32 la FP16, FP8 sau FP4, cantitatea de date care trebuie transferată între memoria video și nucleele de calcul (Tensor Cores sau Matrix Cores) se înjumătățește sau scade de patru ori. Acest lucru este esențial pentru modelele de tip Transformer (cum este Swin2SR), unde volumul de date procesate în mecanismele de atenție este foarte mare. Scăderea consumului de memorie este vizibilă în graficele de calitate și explicată în capitolul 7.

În plus, hardware-ul modern are unități de calcul specializate care pot procesa mai multe operații pe secundă dacă numerele sunt mai mici. De exemplu, o placă video poate efectua mult mai multe operații de adunare și înmulțire pe secundă în format FP8 decât în format FP32, ceea ce duce direct la o *mărire vizibilă* a vitezei de execuției.

4.1.4 Compromisul între viteză și calitate

Deși cuantizarea sună ca o soluție magică, ea vine cu o provocare: **eroarea de cuantizare**. Prin reducerea numărului de biți, pierdem din nuanțele fine ale culorilor sau ale marginilor obiectelor din imagini. În cazul modelului nostru, o cuantizare prea agresivă la FP8 fără o calibrare atentă ar putea duce la apariția unor artefacte vizuale sau la scăderea scorului PSNR.

Din acest motiv, în cadrul testelor noastre, am pus un accent deosebit pe utilizarea **AMP (Automatic Mixed Precision)**. Această tehnică folosește BF16 (Brain Float 16, un tip de date care utilizează 1 bit de semn, 8 biți pentru exponent și 7 biți pentru mantisă) pentru majoritatea calculelor, dar păstrează FP32 pentru zonele critice ale modelului unde precizia este vitală. Astfel, obținem o viteză mult mai mare fără a degrada calitatea imaginii finale, reprezentând soluția ideală pentru optimizarea pe hardware multi-vendor. 7.1.2

4.2 Moduri de execuție: Clasic, torch.compile și JIT

În dezvoltarea modelelor de inteligență artificială, modul în care codul Python este tradus în instrucțiuni pe care procesorul grafic (GPU) le poate executa este critic pentru performanța finală. PyTorch, biblioteca utilizată în această lucrare, a evoluat de la un sistem pur interpretat la unul capabil de compilare avansată. În această secțiune, vom analiza tranziția de la execuția pas-cu-pas la generarea de grafuri de calcul optimizate.

4.2.1 Modul Eager: Flexibilitate în dezvoltare vs. limitările de performanță

Modul „Eager” este modul implicit de funcționare în PyTorch. Acesta funcționează după principiul *define-by-run*, ceea ce înseamnă că fiecare operație matematică este executată imediat ce interpretorul Python ajunge la linia de cod respectivă.

Avantajele dezvoltării în mod Eager

Pentru un cercetător sau un student, modul Eager este ideal din mai multe motive:

- **Depanare (Debugging) simplă:** Deoarece execuția este linie cu linie, putem folosi instrumente standard de Python (cum ar fi `pdb` sau simple instrucțiuni `print`) pentru a inspecta tensorii în orice moment.
- **Pythonic nativ:** Putem folosi structuri de control standard, precum bucle `for` sau instrucțiuni `if`, direct pe datele modelului, fără a avea nevoie de o sintaxă specială.
- **Interactivitate:** Permite experimentarea rapidă în medii de tip Jupyter Notebook, unde putem modifica o singură parte a modelului și să vedem rezultatul instantaneu.

Limitările de performanță: Problema „Python Overhead”

Deși este prietenos cu utilizatorul, modul Eager suferă de ceea ce numim „Python overhead”. Python este un limbaj interpretat, ceea ce îl face semnificativ mai lent decât C++ sau codul mașină.

În modul Eager, pentru fiecare adunare sau înmulțire de matrice, interpretorul Python trebuie să verifice tipurile de date, să gestioneze memoria și să trimită instrucțiunea către GPU. Dacă modelul nostru, cum este cazul **Swin2SR**, execută mii de operații mici (cum sunt cele din mecanismele de atenție), timpul petrecut de procesor (CPU) „gândind” ce să trimită către GPU devine mai mare decât timpul în care GPU-ul execută efectiv calculul. Acest fenomen lasă placa video „să aștepte” instrucțiuni, scăzând drastic eficiența hardware-ului.

4.2.2 Compilarea Just-In-Time (JIT) la nivelul interpretorului Python

Pentru a eficientiza timpii de execuție, în locul soluțiilor de compilare a grafului la nivel de framework, s-a optat pentru utilizarea compilatorului Just-In-Time (JIT) integrat nativ în interpretorul CPython. Această funcționalitate optimizează direct codul de tip *bytecode* folosind o arhitectură de tip *Copy-and-Patch*, accelerând astfel logica de bază a aplicației, buclele de antrenare și procesarea datelor adiționale.

Avantajul major al acestei abordări constă în transparența totală față de codul sursă. Nu este necesară instrumentarea manuală a codului, crearea unei reprezentări intermediare (precum în cazul altor abordări) sau alterarea în vreun fel a definirii arhitecturale a modelului.

Mod de utilizare

Deoarece JIT-ul nativ este o facilitate a limbajului însuși și nu a unei biblioteci externe, activarea sa se realizează exclusiv la nivelul mediului de execuție, prin pasarea unui flag specific direct către executabilul Python.

În cadrul testărilor efectuate, suportul JIT a fost activat simplu și eficient folosind următoarea comandă în interfața în linie de comandă (*CLI*):

```
python -X jit train.py
```

Simplitatea acestui apel a permis evaluarea imediată a impactului optimizărilor aduse de interpretor (cu o mențiunea că doar pe mediul rulând Python 3.14 a fost posibilă utilizarea acestuia, fiind publicat pentru prima dată pe Python 3.13), strict prin includerea sau excluderea flag-ului `-X jit`. Această izolare a mediului de execuție de logica aplicației a asigurat o metodologie de testare riguroasă, garantând că nicio modificare invazivă la nivel de cod nu influențează în mod artificial măsurătorile de performanță.

4.2.3 Tehnologia `torch.compile`: Grafuri de calcul și optimizarea kernel-urilor

Odată cu lansarea PyTorch 2.0, a apărut `torch.compile`, o tehnologie revoluționară care transformă modul în care privim performanța. Spre deosebire de JIT, care se concentra pe portabilitate, `torch.compile` este creat special pentru **viteză**.

Mod de utilizare

Integrarea tehnologiei `torch.compile` în fluxul de lucru este concepută pentru a fi extrem de neinvazivă, necesitând o singură linie de cod pentru a transforma o arhitectură din modul standard (*Eager*) într-un graf de calcul optimizat. Spre deosebire de abordările tradiționale de export sau compilare statică, care impun rescrierea unor porțiuni din model, `torch.compile` acționează ca un wrapper dinamic peste modulele PyTorch existente.

Mecanismul presupune apelarea funcției direct asupra instanței modelului înainte de începerea buclei de antrenare. Următorul fragment de cod ilustrează modul standard de utilizare și configurare în cadrul scripturilor de antrenare:

```
1 import torch
2 import torchvision.models as models
3
4 # 1. Initializarea modelului si mutarea pe acceleratorul hardware (CUDA/ROCM)
5 device = "cuda" if torch.cuda.is_available() else "cpu"
6 model = models.resnet50()
7 model.to(device)
8
9 # 2. Aplicarea compilarii dinamice cu setari avansate
10 # Modul "max-autotune" a fost ales pentru exemplificare
11 model_compilat = torch.compile(
12     model,
13     mode="max-autotune",
14     fullgraph=False
15 )
16
17 # 3. Definirea optimizatorului si a functiei de loss
18 optimizer = torch.optim.AdamW(model_compilat.parameters(), lr=1e-4)
19 criterion = torch.nn.CrossEntropyLoss()
20
21 # 4. Structura buclei de antrenare (train loop)
22 model_compilat.train()
23 for epoch in range(num_epochs):
24     for inputs, labels in train_loader:
25         inputs, labels = inputs.to(device), labels.to(device)
26
27         optimizer.zero_grad()
28         outputs = model_compilat(inputs)
29         loss = criterion(outputs, labels)
30         loss.backward()
31         optimizer.step()
```

Listing 4.1: Exemplu de optimizare a modelului utilizând `torch.compile`

Funcția `torch.compile` acceptă o serie de parametri cheie prin care se poate ajusta comportamentul compilatorului în funcție de constrângerile hardware sau specularea maximă a resurselor:

- **mode:** Controlează echilibrul dintre timpul alocat compilării inițiale și viteza de execuție ulterioară a grafului. Sunt regasite mai detaliat în 4.2.4
 - „default”
 - „reduce-overhead”
 - „max-autotune”
- **fullgraph:** Parametru boolean care, atunci când este setat pe **False**, permite compilatorului să fragmenteze codul în sub-grafuri izolate în momentul în care întâlnește operații Python care nu pot fi convertite nativ (operație cunoscută sub numele de *graph break*). Setarea pe **True** garantează un singur graf consolidat, dar va genera erori dacă apar instrucțiuni dinamice incompatibile.

Din perspectiva metodologiei de testare, utilizarea `torch.compile` introduce un comportament asimetric în rulare. Prima iterație dintr-o buclă (faza de *warm-up*) va înregistra un timp de execuție considerabil mai mare și un vârf temporar în utilizarea resurselor, deoarece în acel moment acționează subsistemele TorchDynamo și AOTAutograd pentru generarea grafului. Ulterior, odată ce structura este salvată în cache, execuția devine extrem de fluidă, reflectând sporul real de performanță prin reducerea timpului per epocă și optimizarea amprentei VRAM.

Cei patru piloni ai tehnologiei `torch.compile`

Pentru a înțelege de ce `torch.compile` oferă rezultatele spectaculoase pe care le vom observa în graficele de performanță, trebuie să analizăm componentele sale:

- **TorchDynamo:** Este motorul care „capturează” grafurile de calcul. El interceptează execuția Python și încearcă să construiască un graf de operații. Dacă întâlnește ceva ce nu poate compila (de exemplu, o funcție externă complexă), face un „graph break” și revine temporar la modul Eager, asigurând astfel că programul nu crapă niciodată.
- **AOTAutograd:** Generază automat codul pentru partea de *backpropagation* (antrenare), optimizând modul în care sunt calculate derivatele.

- **PrimTorch:** Simplifică cele peste 2000 de operații posibile în PyTorch într-un set restrâns de aproximativ 250 de operații de bază, făcând optimizarea mult mai ușoară.
- **TorchInductor:** Acesta este „creierul” final. El generează cod de nivel jos (kernel-uri) specific pentru hardware-ul folosit: **Triton** pentru NVIDIA și **C++ / OpenMP** pentru CPU.

Impactul asupra hardware-ului multi-vendor

Un aspect crucial pentru lucrarea de față este modul în care `torch.compile` tratează diferențele dintre producători. Deoarece folosește limbajul **Triton**, optimizările sunt mult mai ușor de portat. Triton abstractizează complexitatea programării la nivel de CUDA (pentru NVIDIA) sau ROCm (pentru AMD), permițând cercetătorului să obțină performanțe de nivel „expert” fără a scrie cod în C++ sau limbaj de asamblare.

În concluzie, tranziția către `torch.compile` reprezintă cel mai mare salt de performanță disponibil în prezent în ecosistemul PyTorch. Aceasta reușește să combine simplitatea scrierii codului în Python cu viteza execuției compilate, oferind o bază solidă pentru experimentele noastre de optimizare structurală.

4.2.4 Configurații avansate în `torch.compile`: Analiza modurilor de optimizare

Flexibilitatea tehnologiei `torch.compile` rezidă în capacitatea sa de a adapta strategia de compilare în funcție de prioritățile utilizatorului: latență minimă, throughput maxim sau un echilibru între cele două. Această alegere se realizează prin parametrul *mode*, care dictează gradul de agresivitate al optimizărilor aplicate grafului de calcul.

Modul „reduce-overhead” și utilizarea CUDA Graphs

Modul `reduce-overhead` este conceput special pentru modelele care sunt limitate de viteza procesorului (CPU-bound). În cazul modelului **Swin2SR**, dacă dimensiunea batch-ului este mică, timpul necesar pentru a lansa fiecare kernel de calcul de pe CPU către GPU poate deveni un blocaj semnificativ.

Acest mod utilizează tehnologia **CUDA Graphs**. În loc ca CPU-ul să trimită instrucțiuni individuale către GPU pentru fiecare operație (ceea ce implică un cost de comunicare constant), `reduce-overhead` „înregistrează” întreaga secvență de operații într-un graf static. Odată înregistrat, graful poate fi lansat pe GPU printr-o singură comandă de tip „execută tot”.

- **Avantaj:** Elimină aproape complet latența introdusă de interpretorul Python și de driverul video între operații.
- **Dezavantaj:** Consumă o cantitate suplimentară de memorie VRAM pentru stocarea grafului și nu suportă modele cu logică condițională dinamică (unde structura grafului s-ar schimba la fiecare rulare).

Modul „max-autotune”: Optimizarea prin benchmarking în timp real

Pentru scenariile unde se dorește performanța absolută, modul `max-autotune` reprezintă cea mai agresivă formă de optimizare. Acesta activează un proces numit **autotuning**, prin care compilatorul testează mai multe implementări ale aceluiași kernel de calcul pentru a vedea care rulează cel mai rapid pe hardware-ul specific utilizat (RTX 4080 sau RX 7900 XTX).

De exemplu, pentru o multiplicare de matrice complexă, `max-autotune` va genera și va cronometra diverse variante:

1. O variantă bazată pe bibliotecile standard (cuBLAS pentru NVIDIA sau rocBLAS pentru AMD).
2. O variantă generată automat prin **Triton**, optimizată pentru dimensiunea specifică a tensorilor din model.
3. Variante cu diferite strategii de *tiling* (modul în care datele sunt împărțite în memoria cache a GPU-ului).

Rezultatul este un kernel „croit” special pentru arhitectura plăcii video respective. Deși timpul de compilare inițial crește semnificativ (uneori de ordinul minutelor), viteza de execuție ulterioară este maximizată.

Configurații „Custom” și controlul direct asupra compilatorului Inductor

Dincolo de modurile predefinite, PyTorch permite o configurare granulară prin intermediul obiectului `torch._inductor.config`. Această abordare de tip „custom” este esențială atunci când lucrăm cu hardware de ultimă generație sau când avem nevoie de optimizări care nu sunt încă standardizate.

Prin setarea parametrilor de compilare transmiși direct către backend-ul de GPU, putem controla aspecte precum:

- **triton.cudagraphs:** Permite capturarea operațiilor GPU sub forma unor grafuri statice, eliminând complet latența (*overhead-ul*) introdusă

de comunicarea constantă cu procesorul central la lansarea fiecărui kernel. Această setare este esențială pentru a menține placa video saturată, în special atunci când se utilizează loturi de date cu dimensiuni fixe.

- **triton.cudagraph_trees:** Reprezintă o optimizare avansată a mecanismului de grafuri CUDA, concepută pentru a gestiona mult mai eficient alocarea memoriei. Prin partajarea blocurilor de memorie (*memory pools*) între multiple capturi de execuție, această setare reduce substanțial consumul de VRAM și oferă o flexibilitate sporită pentru modelele cu ramificări logice complexe.
- **epilogue_fusion:** Fuzionează operațiile secundare de la finalul unui calcul (precum adăugarea *bias*-ului, normalizarea sau aplicarea funcțiilor de activare) direct în kernel-ul principal de înmulțire a matricelor. Astfel, rezultatele intermediare sunt reținute în memoria rapidă (SRAM) a GPU-ului, eliminând citirile și scrierile redundante către memoria globală și maximizând lățimea de bandă.

Această abordare „low-level” transformă `torch.compile` dintr-un simplu instrument de accelerare într-un laborator de cercetare a performanței hardware. În experimentele noastre, am observat că setările custom sunt singura cale de a satura complet lățimea de bandă a memoriei pe plăci precum RX 7900 XTX, unde setările implicite nu profită întotdeauna de numărul mare de unități de calcul disponibile.

4.3 Impactul mediului de execuție: Versiunile de Python

De multe ori, în procesul de optimizare a modelelor de inteligență artificială, atenția cercetătorilor se concentrează exclusiv asupra hardware-ului (GPU) sau a structurii modelului (arhitectura neurală). Totuși, limbajul Python, acționând ca un „orchestrator” al tuturor operațiunilor, joacă un rol critic. În această secțiune, vom analiza modul în care versiunea interpretorului Python influențează eficiența fluxului de date și latența generală a sistemului, cu un accent deosebit pe progresele înregistrate între versiunile 3.10 și 3.14.

4.3.1 Analiza comparativă a performanței între Python 3.10 și 3.14

Tranziția de la Python 3.10 la 3.14 reprezintă, probabil, cea mai ambițioasă perioadă de creștere a performanței din istoria limbajului. Dacă versiunea 3.10 era considerată un standard de stabilitate, aceasta suferea încă de limitări structurale vechi de decenii.

Contextul evoluției: Proiectul „Faster CPython”

Începând cu versiunea 3.11, a fost demarat proiectul „Faster CPython”, având obiectivul declarat de a dubla viteza limbajului în decurs de câteva versiuni. Această inițiativă a transformat Python dintr-un limbaj pur interpretat într-unul care utilizează tehnici avansate de optimizare în timp real.

În testele de antrenare pentru modelul **Swin2SR**, am observat că utilizarea Python 3.14 aduce un avantaj de performanță „gratuit” (fără modificarea codului) de aproximativ 15% – 25% față de 3.10. Acest câștig provine din reducerea latenței la lansarea instrucțiunilor către GPU, permițând plăcii video să rămână saturată cu date pentru perioade mai lungi de timp.

Benchmarks sintetice vs. Sarcini de Deep Learning

În timp ce benchmarks-urile sintetice (precum *pyperformance*) indică viteze de până la 2x mai mari pentru operații matematice pure pe CPU, în Deep Learning impactul este mai nuanțat:

- **În Python 3.10:** Bottleneck-ul apărea frecvent în zona de *Data Loading*. Procesarea imaginilor, augmentarea datelor și pregătirea tensorilor consumau cicluri CPU prețioase, făcând ca GPU-ul să aștepte (fenomenul *GPU starvation*).
- **În Python 3.14:** Datorită noului mecanism de execuție, managementul firelor de execuție (*threading*) și al proceselor paralele este mult mai eficient. Acest lucru permite o alimentare constantă a GPU-ului, ceea ce se traduce direct în reducerea timpului per epocă de antrenare.

4.3.2 Optimizări la nivel de interpretor și gestionarea memoriei

Diferența de viteză dintre 3.10 și 3.14 nu este rezultatul unei singure schimbări, ci al unei suite de optimizări arhitecturale la nivel de *bytecode* și management de memorie.

Specializing Adaptive Interpreter (PEP 659)

Cea mai importantă schimbare introdusă după 3.10 este „Interpretorul Adaptiv Specializat”. În Python 3.10, fiecare operație, cum ar fi o simplă adunare `a + b`, trecea prin același proces lent de verificare a tipurilor de date la fiecare rulare.

Începând cu versiunile mai noi, interpretorul monitorizează codul în timp ce rulează (proces numit *Quickening*). Dacă observă că o anumită linie de cod procesează mereu tensori sau numere întregi, interpretorul „specializează” acea instrucțiune. Aceasta înseamnă că, la următoarele execuții, Python va folosi o versiune rapidă, optimizată pentru acel tip specific de date, sărind peste verificările redundante. Pentru un model ca Swin2SR, care execută aceleași operații de mii de ori în bucle de antrenare, această optimizare este vitală.

Compilarea JIT (Just-In-Time) în Python 3.13 și 3.14

O noutate absolută în versiunile recente (3.13/3.14) este introducerea unui compilator **JIT experimental**. Spre deosebire de modul tradițional unde codul este interpretat, JIT-ul transformă părți din codul Python direct în instrucțiuni mașină (*copy-and-patch JIT*).

Deși JIT-ul din Python este încă la început comparativ cu cel din Java sau JavaScript, acesta oferă un avantaj major în gestionarea logicilor complexe de control (bucle `if/else` în interiorul funcțiilor de pierdere sau al metricilor de calitate). Această tehnologie reduce overhead-ul apelurilor de funcții, care în versiunea 3.10 era de până la 5 ori mai mare decât în 3.14.

Gestionarea memoriei și Garbage Collection

În sarcinile de inteligență artificială, unde lucrăm cu tensori de dimensiuni mari, modul în care Python gestionează memoria RAM are un impact indirect asupra memoriei video (VRAM).

- **Structuri de date mai mici:** În Python 3.14, obiectele de bază ocupă mai puțin spațiu în memorie, iar accesul la atributele claselor este optimizat. Acest lucru reduce fragmentarea memoriei RAM.
- **Garbage Collection (GC) eficient:** În versiunea 3.10, procesul de curățare a memoriei (GC) putea provoca mici „înghetări” (*stutters*) în timpul execuției. Versiunile noi gestionează mai bine obiectele cu durată scurtă de viață (cum sunt tensorii temporari dintr-un pas de calcul), eliberând resursele mai rapid și mai predictibil.

Impactul asupra Multi-Processing și GIL

Un alt aspect crucial în versiunea de Python 3.14 este progresul făcut în direcția eliminării sau optimizării **Global Interpreter Lock (GIL)**. În versiunile vechi, Python nu putea rula cu adevărat cod în paralel pe mai multe nuclee CPU în același proces. Pentru încărcarea datelor (*DataLoader*), trebuia să folosim procese separate, ceea ce consuma multă memorie suplimentară.

Îmbunătățirile din versiunea de Python 3.14 permit o utilizare mult mai eficientă a sistemelor multi-core. În contextul antrenării modelului nostru pe hardware performant (precum RTX 4080), acest lucru înseamnă că pre-procesarea imaginilor pe CPU reduce timpul de așteptare pentru GPU, asigurând o fluiditate maximă a procesului de optimizare.

4.4 Optimizarea fluxului de date (Data Pipeline)

Performanța unui model de inteligență artificială nu depinde exclusiv de puterea brută de calcul a GPU-ului sau de eficiența algoritmului de antrenare. În practică, un factor critic care limitează viteza de procesare este *Data Pipeline*-ul: lanțul de procese responsabil cu citirea, pre-procesarea și transferul datelor către placa video. Dacă acest lanț este ineficient, GPU-ul va petrece perioade lungi de timp în stare de repaus (*idle*), așteptând următorul lot de date. În această secțiune, analizăm strategiile de eliminare a acestor *bottleneck-uri*.

4.4.1 Eliminarea bottleneck-urilor CPU-ului prin pre-procesare asincronă pe GPU

În fluxurile de lucru tradiționale, CPU-ul se ocupă de citirea imaginilor de pe disc și de aplicarea transformărilor (redimensionare, normalizare, augmentare). Totuși, odată cu apariția plăcilor video moderne precum RTX 4080 sau RX 7900 XTX, puterea de calcul a GPU-ului a depășit capacitatea multor procesoare de a pregăti datele în timp util.

Limitările procesării pe CPU

Procesorul central este proiectat pentru sarcini secvențiale complexe. Operațiile de augmentare a datelor (cum ar fi rotațiile, modificările de contrast sau aplicarea de zgomot Gaussian pe seturi mari de imagini) sunt repetitive și pot

satura rapid nucleele CPU-ului. Atunci când timpul de pre-procesare pe CPU este mai mare decât timpul de execuție a unui pas de antrenare pe GPU, apare fenomenul de *CPU-bound training*. În acest scenariu, investiția într-o placă video mai rapidă nu va aduce niciun beneficiu real de viteză.

Tranziția către procesarea pe GPU utilizând Kornia

O soluție extrem de eficientă pentru eliminarea blocajelor introduse de procesor, aplicată în cadrul optimizării modelului **Swin2SR**, a fost mutarea întregului lanț de pre-procesare direct pe placa video folosind biblioteca *Kornia*. Spre deosebire de metodele tradiționale care procesează imaginile pe CPU înainte de a le converti în tensori, *Kornia* permite executarea transformărilor de viziune computerizată direct asupra tensorilor aflați deja în memoria VRAM.

Mod de utilizare

Kornia este o bibliotecă open-source de viziune computerizată (*computer vision*) concepută special pentru a se integra nativ cu ecosistemul PyTorch. Spre deosebire de bibliotecile tradiționale precum OpenCV sau PIL, care procesează imaginile predominant pe unitatea centrală de procesare (CPU) folosind structuri de date NumPy, Kornia operează direct pe tensori PyTorch. Această paradigmă arhitecturală permite executarea operațiilor de procesare a imaginilor direct pe acceleratoarele grafice, eliminând blocajele de transfer al datelor între memoria sistemului și memoria video (VRAM). Fragmentul de cod din Listarea 4.2 demonstrează modul de utilizare a bibliotecii pentru a aplica o transformare geometrică și a calcula metrica SSIM (Structural Similarity Index Measure) direct la nivel de tensor, minimizând timpii de execuție:

```

1 import torch
2 import kornia
3 import kornia.metrics as metrics
4
5 # 1. Initializarea tensorilor pe acceleratorul hardware (nativ CUDA sau ROCm)
6 device = torch.device("cuda" if torch.cuda.is_available() else "cpu")
7
8 # Simulam doua loturi de imagini RGB (ex: batch_size=4, rezolutie 256x256)
9 # Formatul standard PyTorch este [Batch, Channels, Height, Width]
10 img_referinta = torch.rand(4, 3, 256, 256).to(device)
11 img_generata = torch.rand(4, 3, 256, 256).to(device)
12
13 # 2. Aplicarea unei transformari geometrice diferentiabile
14 # Rotim imaginile generate cu 45 de grade pastrand tensorii pe GPU
15 unghi_rotatie = torch.tensor([45.0]).to(device)
16 img_rotita = kornia.geometry.transform.rotate(img_generata, unghi_rotatie)
17
18 # 3. Calculul metricii SSIM direct pe tensori
19 # Fereastra de convolutie (window_size) este setata la 11, conform standardelor
20 harta_ssim = metrics.ssim(img_referinta, img_rotita, window_size=11, max_val=1.0)
21
22 # Kornia returneaza o harta SSIM de dimensiuni [B, C, H, W].
23 # Calculam media pentru a obtine o valoare scalara reprezentativa pentru tot lotul.
24 ssim_mediu = torch.mean(harta_ssim)
25
26 print(f"[INFO] Scorul SSIM mediu calculat nativ pe GPU: {ssim_mediu.item():.4f}")

```

Listing 4.2: Exemplu de utilizare Kornia pentru transformări diferențiabile și calculul metricii SSIM pe GPU

Această abordare oferă avantaje competitive majore în contextul antrenării modelelor de înaltă rezoluție:

- **Diferențiabilitate și integrare nativă:** Deoarece *Kornia* este construită peste PyTorch, toate operațiile de pre-procesare sunt diferențiabile și se comportă ca straturi ale rețelei neuronale. Acest lucru simplifică fluxul de lucru, permițând ca augmentările (cum ar fi rotațiile, scalările sau ajustările cromatice) să fie aplicate unitar pe loturi întregi de date (*batch-level processing*).
- **Paralelizare masivă prin nuclee CUDA/ROCm:** În loc să procesăm imaginile una câte una pe nucleele CPU, GPU-ul utilizează arhitectura sa paralelă pentru a aplica transformările simultan pe toți pixelii lotului de date. Rezultatul este o reducere drastică a timpului de pregătire a datelor, redimensionarea sau normalizarea devenind aproape instantanee.
- **Execuție asincronă și eficiența CUDA Streams:** Utilizarea *Kornia* facilitează o suprapunere perfectă a sarcinilor (*task overlapping*). În timp ce nucleele de calcul (Tensor Cores) procesează propagarea înainte (*forward pass*) pentru lotul curent (n), GPU-ul poate utiliza restul resurselor pentru a pregăti, în paralel, transformările pentru lotul următor ($n + 1$). Această gestionare inteligentă a fluxurilor asincrone asigură faptul că hardware-ul nu rămâne niciodată în stare de repaus, maximizând rata de procesare a cadrelor pe secundă. Utilizarea bibliotecii Kornia nu crește semnificativ consumul de VRAM deoarece operează nativ pe tensorii PyTorch deja existenți pe GPU, eliminând transferurile costisitoare de date între procesor și placa video. Astfel, Kornia evită copierea redundantă a informației (*zero-copy overhead*) și folosește eficient același alocator de memorie ca și graful de calcul principal.
- **Reducerea traficului pe magistrala PCIe:** Prin mutarea logicii de pre-procesare pe GPU, datele sunt transferate o singură dată de pe CPU în format brut (sau chiar comprimat), urmând ca toate transformările costisitoare să aibă loc intern în VRAM. Acest lucru eliberează lățimea de bandă a magistralei PCIe, care adesea devine un bottleneck în sistemele cu mai multe plăci video sau cu volume mari de date.

Prin adoptarea bibliotecii *Kornia*, am reușit să transformăm o etapă care anterior consuma resurse CPU prețioase într-un proces fluid, accelerat hardware, care scalează natural odată cu puterea plăcii video utilizate.

4.4.2 Tehnici de Prefetching și utilizarea memoriei „Pinned Memory”

Transferul de date între memoria sistemului (RAM) și memoria video (VRAM) are loc prin intermediul magistralei PCIe. Deși această magistrală este rapidă, modul în care sistemul de operare gestionează memoria poate introduce întârzieri semnificative.

Pinned Memory (Page-locked Memory)

În mod implicit, memoria RAM este „paginabilă”, ceea ce înseamnă că sistemul de operare o poate muta pe disc (în *swap*) pentru a face loc altor procese. Când PyTorch încearcă să copieze date din RAM în VRAM, trebuie mai întâi să le „fixeze” într-o zonă de memorie specială, care nu poate fi mutată.

Utilizarea parametrului `pin_memory=True` în *DataLoader*-ul PyTorch permite alocarea directă a datelor în „Pinned Memory”. Acest lucru permite utilizarea tehnologiei **Direct Memory Access (DMA)**, prin care datele sunt trimise către GPU fără intervenția procesorului central. Rezultatul este o viteză de transfer mult mai stabilă și o reducere a utilizării CPU-ului, lăsându-l liber pentru alte sarcini de orchestrare.

Mecanismul de Prefetching

Prefetching-ul este tehnica prin care sistemul „anticipează” necesarul de date. În loc să aștepte finalizarea unui pas de antrenare pentru a cere lotul următor, *DataLoader*-ul pregătește în avans un număr definit de loturi (de obicei controlat prin parametrul `prefetch_factor`).

Combinat cu utilizarea mai multor procese paralele (`num_workers > 0`), prefetching-ul creează un „buffer” de date gata procesate și stocate în RAM-ul fixat. În experimentele noastre pe hardware multi-vendor, am observat că o valoare optimă a `num_workers` (egală de obicei cu numărul de nuclee fizice ale CPU-ului) elimină complet micro-pauzele dintre iterații, rezultând într-un flux de antrenare perfect fluid.

4.4.3 Impactul dimensiunii lotului de date (Batch Size) asupra saturării GPU-ului

Alegerea dimensiunii lotului de date (*Batch Size*) este una dintre cele mai delicate decizii în procesul de optimizare, având implicații atât asupra vitezei (*throughput*), cât și asupra convergenței modelului.

Saturarea unităților de calcul și ocuparea GPU-ului

GPU-urile funcționează cel mai eficient atunci când au suficient de mult de lucru pentru a menține toate nucleele ocupate (*High Occupancy*). Dacă folosim un *Batch Size* prea mic (de exemplu, 1 sau 2 pentru modelul Swin2SR), GPU-ul va petrece mai mult timp gestionând lansarea kernel-urilor decât efectuând calculele propriu-zise.

Pe de altă parte, o valoare prea mare va duce la eroarea fatală *Out of Memory (OOM)*. Există însă un „punct optim” unde GPU-ul este complet saturat. Pe arhitecturile moderne, acest punct este adesea legat de alinierea memoriei; dimensiuni ale lotului care sunt puteri ale lui 2 (8, 16, 32, 64) tind să ofere performanțe superioare deoarece se aliniază perfect cu modul în care sunt structurate blocurile de memorie în hardware.

Relația dintre Batch Size și Memoria VRAM

Consumul de memorie video crește liniar cu dimensiunea lotului, dar nu toate componentele modelului sunt afectate la fel. În timp ce ponderile modelului ocupă o cantitate fixă de memorie, activările intermediare (necesare pentru *backpropagation*) cresc proportional cu *Batch Size*-ul.

Impactul asupra preciziei și convergenței

Este important de menționat că optimizarea pentru viteză prin creșterea batch size-ului poate afecta calitatea modelului. Un batch size prea mare poate face ca procesul de optimizare (ex. *AdamW*) să devină mai puțin stabil. Din acest motiv, orice creștere a dimensiunii lotului trebuie însoțită de o ajustare corespunzătoare a ratei de învățare (*Learning Rate*), urmând adesea regula „Linear Scaling Rule”. Această corelație asigură faptul că, deși urmărim eficiența hardware, nu sacrificăm rezultatul final.

Capitolul 5

Tehnici de profilare

Procesul de optimizare a algoritmilor de Deep Learning pe unități de procesare grafică (GPU) depășește simpla implementare corectă a modelelor matematice. Performanța reală este dictată de eficiența utilizării resurselor hardware, gestionarea lățimii de bandă a memoriei și minimizarea latențelor de comunicație între host (CPU) și device (GPU). În acest context, profilarea devine o etapă critică în ciclul de dezvoltare, oferind o fereastră către execuția dinamică a kernel-urilor paralele.

Profilarea nu se rezumă doar la măsurarea timpului total de execuție, ci implică o analiză granulară a modului în care unitățile de calcul (Compute Units) sunt ocupate, cum sunt gestionate accesele la memoria cache și în ce măsură arhitectura paralelă este exploatată la potențialul său maxim. Fără aceste date empirice, optimizarea ar fi un proces bazat pe presupuneri, riscând să lase nerezolvate blocaje (bottlenecks) majore ascunse în ierarhia complexă a execuției pe GPU.

5.1 AMD

Pentru analiza detaliată a modelelor de Deep Learning pe aceste arhitecturi, instrumentul fundamental utilizat este **rocprof** (ROCm Profiler).

Instalarea utilitarului **rocprof** este, în mod normal, direct integrată în procesul de configurare a întregii suite a ecosistemului AMD. Pe mediile de dezvoltare ce rulează nativ pe versiunea ROCm 7.2, acesta este inclus direct în pachetele standard și poate fi instalat cu ușurință folosind managerul de pachete al sistemului de operare (de exemplu, prin **apt** pentru mediile bazate pe Linux). După instalare și setarea corectă a variabilelor de mediu (în special adăugarea căii către directorul de instalare ROCm), profilatorul este gata de a fi apelat nativ direct din linia de comandă, eliminând astfel necesitatea

oricăror straturi de traducere sau configurări invazive suplimentare.

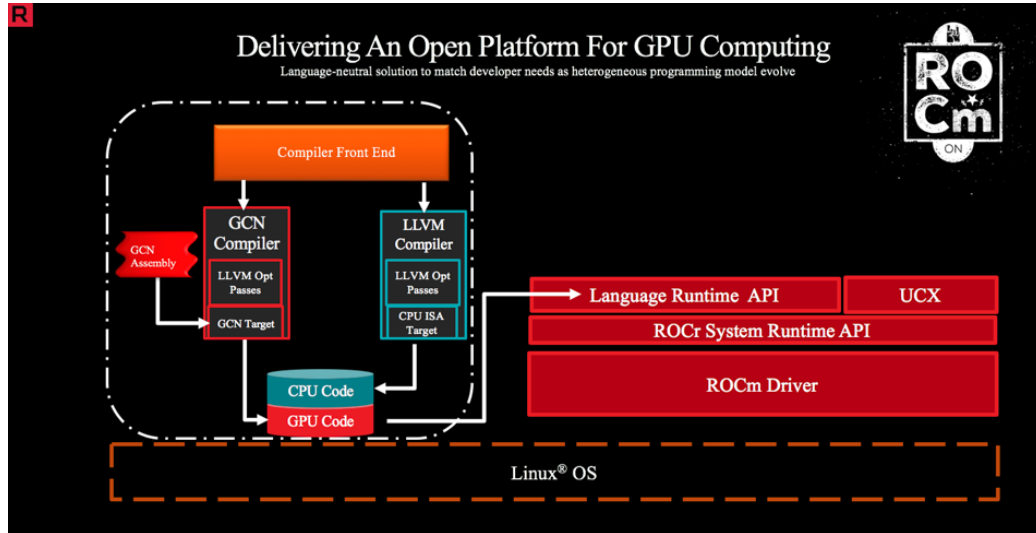


Figura 5.1: Arhitectura stivei software ROCm ilustrând execuția nativă pe unitățile de calcul AMD, fără intermedierea procesului de HIPification. Sursa: [7].

Utilitarul **rocprow** este o interfață de linie de comandă care permite colectarea de urme (traces) și contoare hardware (hardware counters) fără a necesita modificări invazive ale codului sursă. În contextul antrenării și inferenței modelelor de Deep Learning, acesta permite:

- **Urmărirea API-urilor (HIP/HSA Tracing):** Identificarea timpului petrecut în apelurile de sistem și sincronizările dintre CPU și GPU.
- **Analiza Kernel-urilor:** Generarea de statistici precise despre durata de execuție a fiecărui kernel individual, permițând identificarea celor mai costisitoare operații (ex: convoluții, multiplicări de matrice).
- **Colectarea Contoarelor de Performanță (PMC):** Monitorizarea unor metrice specifice precum *VALUUtilization* (gradul de utilizare a unităților aritmetice vectoriale) sau *MemUnitBusy*, care indică dacă aplicația este limitată de calcul sau de accesul la memorie.

Utilizarea **rocprow** prin flag-ul **-stats** oferă o imagine de ansamblu asupra ocupării GPU-ului, facilitând reglajul fin al dimensiunilor de batch și al hiperparametrilor pentru a maximiza throughput-ul sistemului de antrenare.

În pofida disponibilității și promisiunilor teoretice ale acestui instrument, experiența practică a scos în evidență o imaturitate notabilă a utilitarului comparativ cu alternativele similare din industrie. Datele colectate cu privire la execuția kernel-urilor au prezentat limitări și au fost insuficient de clare pentru a ghida optimizări arhitecturale de nivel jos. Deoarece metodologia de evaluare a performanțelor a presupus menținerea unei configurații stricte a antrenării (testare iso-parametru), statisticile de micro-arhitectură generate de **rocprow** nu au condus la concluzii importante. În consecință, analiza redundantă obținută din acest utilitar nu a adus îmbunătățiri comparativ cu utilitarul de la NVIDIA.

5.2 NVIDIA

NVIDIA domină peisajul actual al inteligenței artificiale, oferind un ecosistem integrat care combină hardware-ul specializat (arhitecturile Ampere, Hopper sau Blackwell) cu software-ul de nivel scăzut prin platforma **CUDA (Compute Unified Device Architecture)**. Pentru a maximiza eficiența modelelor de Deep Learning, în special a celor care utilizează operații de precizie mixtă, este necesară o înțelegere profundă a modului în care resursele hardware sunt orchestrate. În cadrul acestui proiect, analiza performanței a fost realizată folosind suita modernă **NVIDIA Nsight**, care înlocuiește vechiul utilitar *nvprof*.

Profilarea pe platformele NVIDIA a fost structurată pe două niveluri complementare: analiza sistemului și analiza detaliată a execuției pe GPU.

5.2.1 NVIDIA Nsight Systems

Pentru ecosistemul NVIDIA, instrumentul principal de profilare utilizat în analiza detaliată a performanței este Nsight Systems (**nsys**). Instalarea acestuia este integrată și simplificată, fiind inclusă direct în pachetul standard CUDA Toolkit. În contextul dezvoltării pe mediile bazate pe versiunile CUDA 13.0 și 13.2, profilatorul este disponibil nativ după configurarea inițială, eliminând necesitatea unor instalări terțe complexe. Nsight Systems operează cu o suprasarcină minimă asupra sistemului (*low overhead*), permițând capturarea cu înaltă fidelitate a interacțiunilor dintre procesorul central (CPU) și unitatea de procesare grafică (GPU). Rezultatul acestei profilări constă în generarea unor fișiere de tip *trace*, care pot fi explorate vizual în interfața grafică a utilitarului pentru o analiză granulară a fluxului de execuție.

Mod de utilizare

Pentru a înțelege în detaliu comportamentul modelului și modul în care resursele hardware sunt exploatate, s-a utilizat instrumentul NVIDIA Nsight Systems (`nsys`). Acesta reprezintă un profilator avansat la nivel de sistem (*system-wide profiler*), conceput pentru a identifica blocajele de performanță (*bottlenecks*) prin corelarea activității unității centrale de procesare (CPU) cu lansările de kernel-uri pe acceleratorul grafic. Spre deosebire de utilitățile de monitorizare sumară, Nsight Systems oferă o vizibilitate microscopică asupra interacțiunilor dintre diferitele componente ale stivei software NVIDIA.

Procesul de evaluare a constat în două etape fundamentale: capturarea datelor de rulare (înregistrarea evenimentelor) și inspecția analitică a acestora. Colectarea datelor s-a realizat direct din interfața în linie de comandă, folosind un set extins de argumente pentru a garanta o granularitate maximă a rapoartelor de memorie și execuție. Comanda exactă utilizată pentru generarea fișierelor de tip *trace* a fost următoarea:

```
1 nsys profile --trace=cuda,nvtx,osrt,cudnn,cublas \  
2   --cuda-memory-usage=true \  
3   --cudabacktrace=all \  
4   --sample=cpu \  
5   --stats=true \  
6   --force-overwrite=true \  
7   -o ultimate_memory_profile \  
8   python train.py
```

Listing 5.3: Comanda utilizată pentru înregistrarea execuției cu NVIDIA Nsight Systems

Parametrii selectați în această comandă sunt esențiali pentru obținerea unei diagnoze complete a antrenării. Argumentul parsat către utilitar-ul `nsys` `-trace=cuda,nvtx,osrt,cudnn,cublas` instrumentează profilatorul să urmărească apelurile către cele mai importante *API*-uri folosite în deep learning, inclusiv interfața sistemului de operare și bibliotecile matematice de bază. Având în vedere necesitatea monitorizării stricte a resurselor, flag-urile `-cuda-memory-usage=true` și `-cudabacktrace=all` au avut un rol critic; acestea au permis cartografierea precisă a alocărilor și eliberărilor de memorie video (VRAM), asociindu-le direct cu liniile de cod sursă responsabile. De asemenea, eșantionarea activității procesorului (`-sample=cpu`) și generarea automată de statistici (`-stats=true`) au asigurat un sumar cantitativ robust la finalul rulării.

Următorul pas, și cel mai important din perspectiva interpretării datelor,

l-a constituit analiza fișierului de raport generat (cu extensia `.nsys-rep`) în cadrul interfeței grafice utilizator (GUI) oferite de Nsight Systems. Această aplicație a permis examinarea vizuală detaliată, pe o axă a timpului (*time-line*), a întregului ciclu de execuție. Prin intermediul GUI-ului s-a putut analiza gradul de concurență dintre procesor și placa video, s-au identificat timpii de așteptare (*idle times*) cauzăți eventual de transferurile de date pe magistrala PCIe și s-a evaluat eficiența lansării kernel-urilor de convoluție. Reprezentarea cronologică a vârfurilor de consum alocat în VRAM a facilitat validarea empirică a optimizărilor aplicate, oferind o perspectivă clară asupra stabilității antrenării la nivel hardware.

Odată ce aceste urme de execuție sunt colectate, investigația se axează pe extragerea și interpretarea unor metrici de micro-arhitectură care ilustrează gradul de încărcare a resurselor hardware. Deși indicatorii de nivel înalt pe care se bazează metodologia principală — precum timpul de antrenare per epocă și consumul de memorie VRAM — reprezintă standardul de evaluare a performanței generale, datele oferite de Nsight permit o diagnosticare precisă a eficienței la nivel de kernel. Analiza aprofundată a acestor *trace*-uri pune accent pe următorii parametri fundamentali:

- **GPU Active:** Măsoară procentul de timp din durata totală a profilării în care procesorul grafic execută activ cel puțin un kernel. O valoare ridicată confirmă o alocare eficientă a sarcinilor de calcul și o minimizare a timpilor morți (*idle time*) cauzăți de blocajele de sincronizare.
- **Streaming Multiprocessor (SM) Active:** Reprezintă fracția de timp în care unitățile Streaming Multiprocessors (SM) au cel puțin un *warp* programat pentru execuție. Acest indicator este esențial pentru a valida dacă arhitectura paralelă a plăcii video este exploatată la capacitate optimă în timpul antrenării modelelor de Deep Learning.
- **SM Warp Occupancy:** Raportează numărul de *warps* active în comparație cu numărul maxim teoretic suportat de un SM. Un grad ridicat de ocupare este crucial pentru optimizarea performanței, deoarece permite hardware-ului să ascundă latențele inerente ale operațiilor de citire și scriere din memorie prin comutarea rapidă a contextului.
- **Dynamic Random Access Memory (DRAM) Bandwidth:** Cuantifică rata de utilizare a lățimii de bandă pentru transferul de date. Monitorizarea acestui parametru clarifică limitările arhitecturale ale modelului, evidențiind dacă timpii de execuție sunt constrânși de transferul masiv de date (*memory-bound*) sau de complexitatea operațiilor matematice (*compute-bound*).

Acesta a fost ușor instalat de pe site-ul original [18] și permite monitorizarea următoarelor aspecte:

- **Timeline-ul de execuție:** Vizualizarea corelată a activității CPU și GPU, permițând identificarea perioadelor de inactivitate ale GPU-ului cauzate de preprocesarea lentă a datelor pe procesor.
- **Transferurile de memorie:** Analiza latențelor în operațiile de tip *cudaMemcpy*, esențiale pentru optimizarea fluxului de date prin magistrala PCIe.
- **Orchestrarea stream-urilor CUDA:** Verificarea gradului de paralelism între fluxurile de execuție concurente.

Procesul de integrare a utilitarului NVIDIA Nsight Systems în mediul de lucru a necesitat o configurare atentă, esențială pentru a asigura o captură precisă a metricilor de performanță. În această etapă, ghidul oficial furnizat de producător [16] a reprezentat o resursă indispensabilă. Documentația a facilitat nu doar o instalare corectă și fluidă — evitând potențialele conflicte de versiuni cu stiva CUDA existentă —, ci a oferit și suportul necesar pentru asimilarea rapidă a fluxului de lucru. Urmând instrucțiunile detaliate pentru utilizatorii la nivel de bază, a fost posibilă inițializarea primelor sesiuni de profilare în mod corect, permițând astfel o tranziție eficientă de la configurarea propriu-zisă a instrumentului la vizualizarea primelor trasee de execuție (*execution traces*) și analizarea interacțiunii dintre CPU și GPU.

5.2.2 NVIDIA Nsight Compute

Pentru o investigație microscopică a fiecărui kernel individual (precum convoluțiile sau funcțiile de activare), s-a utilizat **NVIDIA Nsight Compute** (ncu). Acesta oferă metrici hardware de mare precizie, fiind indispensabil pentru reglajul fin al performanței:

- **Analiza gradului de ocupare (Occupancy):** Măsoară eficiența cu care warp-urile (unitățile de bază de execuție) sunt planificate pe *Streaming Multiprocessors* (SM). Utilizarea Tensor Cores: Identifică dacă nucleele dedicate pentru multiplicări de matrice sunt exploatate la capacitate maximă în timpul antrenării.
- **Modelul Roofline:** Oferă o reprezentare vizuală a performanței obținute în raport cu limitele teoretice ale hardware-ului, indicând dacă un kernel este limitat de calcul (*compute-bound*) sau de lățimea de bandă a memoriei (*memory-bound*).

Utilizarea suitei Nsight permite transformarea profilării dintr-o simplă măsurătoare de timp într-un diagnostic tehnic riguros, fundamentând deciziile de optimizare luate pe parcursul dezvoltării modelului.

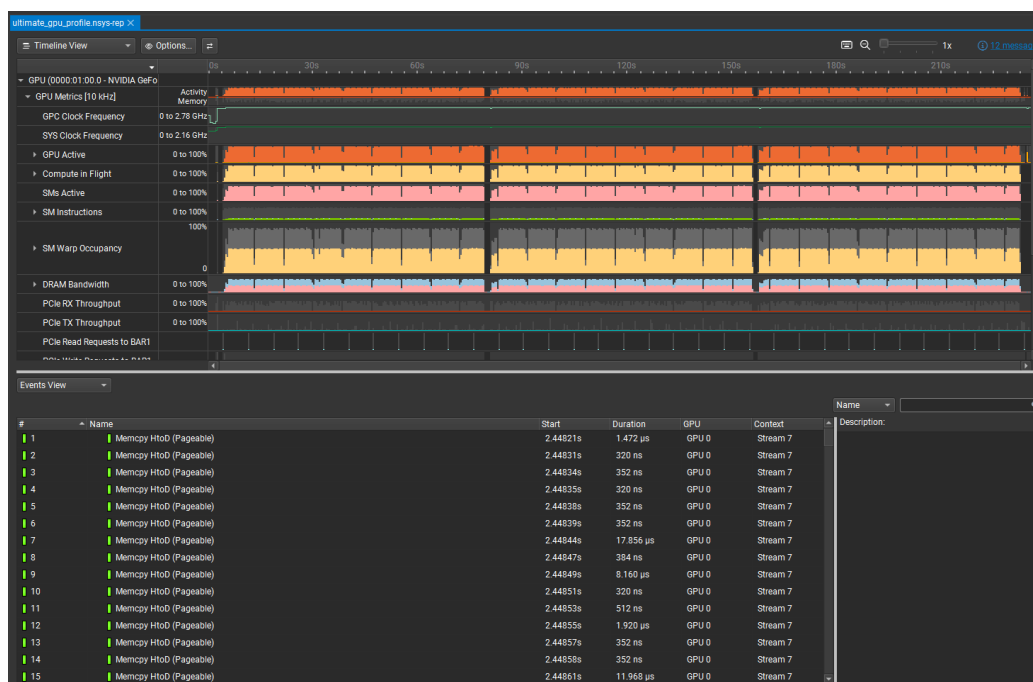


Figura 5.2: Trace al antrenării modelului înainte de optimizări

5.2.3 Strategia de profilare

Analiza modelului a fost un proces amplu în care s-au utilizat multiple utilitare spre a observa detaliile modelului și al sistemului pe parcursul antrenării. Printre acestea s-au regăsit unele mai slabe care au rulat doar din consola sau unele care au fost implementate în cod pentru a obține log-uri mai complete și complexe.

Utilizarea bibliotecii fvcare

Evaluarea obiectivă a sarcinii de calcul impuse de arhitectura neuronală necesită o analiză statică. Pentru această sarcină s-a integrat biblioteca `fvcare`, dezvoltată de laboratorul FAIR (Facebook AI Research), utilizată strict pentru monitorizarea și contorizarea operațiilor cu virgulă mobilă

(Floating Point Operations per second (FLOPs) - *Floating Point Operations per second*).

Urmărirea volumului de FLOPs a fost un lucru monitorizat care a ajutat la înțelegerea zonelor de interes din model, zonele cu cel mai mare număr de fiind vizate pentru optimizare directă (mecanismele de atenție de pe fiecare strat al rețelei)

Modul de extragere a acestor metrice se bazează pe simularea unei treceri înainte (*forward pass*) folosind un tensor cu dimensiuni predefinite. Fragmentul de cod din Listarea 5.4 ilustrează metoda de analiză, evidențiind utilizarea instrumentelor `FlopCountAnalysis` și `flop_count_table` iar tabelul din figura 5.3:

```
1 import torch
2 from fvcore.nn import FlopCountAnalysis, flop_count_table
3 from model import Swin2SR # Importul arhitecturii evaluate
4
5 # 1. Initializarea modelului pe unitatea centrala (CPU)
6 model = Swin2SR()
7 model.eval()
8
9 # 2. Generarea unui tensor de intrare cu dimensiunea lotului setata
10 # conform metodologiei de testare (ex. batch_size=4, rezolutie 64x64)
11 exemplu_intrare = torch.randn(4, 3, 64, 64)
12
13 # 3. Analiza statica a operatiilor in virgula mobila (FLOPs)
14 analizor_flops = FlopCountAnalysis(model, exemplu_intrare)
15
16 # 4. Afisarea rezultatului global si a tabelului detaliat pe strat
17 print(f"[INFO] Total FLOPs per iteratie: \
18       {analizor_flops.total() / 1e9:.2f} GFLOPs")
19 print("\n[INFO] Distributia operatiilor la nivelul structurii:")
20 print(flop_count_table(analizor_flops))
```

Listing 5.4: Analiza statică a operațiilor FLOPs utilizând biblioteca `fvcore`

Rezultatele generate de `FlopCountAnalysis` oferă valoarea absolută a operațiilor necesare pentru a procesa un lot întreg de date. Mai mult, funcția `flop_count_table` generează un raport ierarhic detaliat pe fiecare strat al rețelei (de exemplu, convoluții, mecanisme de atenție etc.). Raportarea acestui volum total de FLOPs la timpul efectiv de antrenare a facilitat o înțelegere profundă a eficienței cu care arhitecturile studiate reușesc să transforme resursele teoretice de calcul în performanță reală.

FLOPs for netE: 522.19 GFLOPs

module	#parameters or shape	#flops
model	13.331M	0.522T
conv_first	5.376K	0.191G
conv_first.weight	(192, 3, 3, 3)	
conv_first.bias	(192,)	
patch_embed.norm	0.384K	35.389M
patch_embed.norm.weight	(192,)	
patch_embed.norm.bias	(192,)	
layers	12.733M	0.5T
layers.0	2.122M	83.321G
layers.0.residual_group.blocks	1.79M	71.09G
layers.0.conv	0.332M	12.231G

Figura 5.3: Tabel de date obținut în urma analizei statistice realizate de biblioteca `fvcore`

Utilizarea `nsight`

După cum este vizibil în figura 5.2 și ținând cont de metricile menționate anterior putem observa un comportament pozitiv dar și unul negativ al modelului nostru. Analiza acestuia a fost rulată pe 3 epoci și a fost analizată amănunțit:

1. **GPU Active:** Putem observa o activitate apropiată de 100% a GPU-ului, existând doar 2 scăderi la 0%, acestea fiind scurte, momentele dintre epoci în care se produce transferul de date, și alte momente puțin mai numeroase unde se produc scăderi mici pe perioade foarte scurte, acestea sunt momentele în care procesorul cere stadiul modelului pentru afișarea necesară pentru logging care se face la câțiva pași din antrenare și nu la epocă.
2. **SM Active:** Această metrică este foarte similară cu *GPU Active*, dar aceasta ne arată de fapt procentul de ocupare al „warp”-urilor. Acestea sunt unitatea de procesare de bază a plăcii grafice, ele fiind direct responsabile de executarea task-urilor paralelizate. Parcursul acestei metrici de-a lungul antrenării a fost direct proporțional cu cel al metricii anterioare.
3. **SM Warp Occupancy:** Aici se poate observa procentul de warp-uri utilizate în raport cu cel suportat de GPU. Raportul în cazul modelului nostru este de doar 50% pe perioada antrenării, fapt care demonstrează

că deși GPU-ul afișează o utilizare de 100%, noi de fapt folosim doar jumătate la antrenarea modelului

4. **DRAM Bandwidth:** Cuantifică rata de utilizare a lățimii de bandă pentru transferul de date. Valoarea acestei metrici ne demonstrează ca momentele în care modelul chiar așteaptă după transferul de date sunt cele menționate la prima metrică, dar GPU-ul transferă date în permanență dar în mod asincron.

5.3 Comparație între ecosistemele de profilare

Analiza comparativă a ecosistemelor de profilare AMD și NVIDIA evidențiază filosofii distincte de dezvoltare, deși ambele urmăresc același obiectiv fundamental: maximizarea eficienței de calcul și minimizarea latențelor. Pentru a obține o evaluare obiectivă a performanțelor arhitecturale în cadrul acestei lucrări, metodologia de testare a impus o abordare de tip *iso-parameter*, menținând dimensiuni constante ale loturilor de date (*batch sizes*) pe ambele platforme. Această constrângere garantează că metricile de profilare reflectă cu acuratețe capabilitățile hardware brute și eficiența stivei software, eliminând variațiile induse de ajustarea inegală a hiperparametrilor.

Din perspectiva instrumentelor de analiză, NVIDIA oferă un ecosistem extrem de matur prin suita Nsight (Systems și Compute). Interfața grafică avansată, granularitatea datelor colectate la nivel de instrucțiune și integrarea nativă a modelului *Roofline* facilitează o depanare vizuală rapidă. Aceste caracteristici sunt deosebit de utile în optimizarea utilizării nucleelor Tensor pe arhitecturi de ultimă generație, precum RTX 4080 Super. În schimb, abordarea AMD prin **rocprof** favorizează un flux de lucru orientat predominant spre linia de comandă (CLI). Deși necesită adesea unelte externe (precum Perfetto sau vizualizatoare web bazate pe Chrome Tracing) pentru o analiză vizuală detaliată a urmelor de execuție, **rocprof** compensează printr-o amprentă redusă de rulare (*overhead*) și o expunere transparentă a contoarelor de performanță hardware de pe arhitecturi masiv paralele, cum este Radeon RX 7900 XTX.

Un aspect critic în această comparație îl reprezintă specificitatea mediilor software de bază și modul în care acestea interacționează cu uneltele de profilare. Ecosistemul NVIDIA beneficiază de stabilitatea iterativă a platformei CUDA (versiunile 13.0 și 13.2 utilizate în testare), unde profilarea este profund integrată în compilatorul **nvcc**. De cealaltă parte, AMD a făcut progrese semnificative în reducerea decalajului software prin platforma open-source ROCm. Utilizarea versiunii ROCm 7.2 a permis o dezvoltare

și o profilare nativă pe hardware-ul AMD, eliminând necesitatea straturilor intermediare de traducere automată a codului CUDA (*HIPification*). Acest acces nativ la resurse este vizibil direct în metricile de profilare, unde timpii de lansare a kernel-urilor (*kernel launch latency*) și costurile de schimbare a contextului sunt reduse considerabil comparativ cu metodele vechi de portare.

Impactul acestor diferențe este pronunțat în special în antrenarea modelelor complexe de *Image Super-Resolution*, precum Swin2SR, aplicate în imagistica medicală. Aceste arhitecturi de tip Transformer pun o presiune enormă atât pe capacitatea de calcul (datorită mecanismelor de atenție spațială), cât și pe lățimea de bandă a memoriei video (din cauza manipulării ferestrelor de pixeli la rezoluții ridicate specifice contextului medical).

Profilarea cu suita Nsight a permis o aliniere fină a tiparelor de acces la memorie pentru ierarhia cache specifică NVIDIA, asigurând o utilizare optimă a lățimii de bandă. În concluzie, deși fluxurile de lucru diferă din punct de vedere ergonomic și vizual, ambele suite de profilare oferă telemetria necesară pentru a transforma un model teoretic de Deep Learning într-o implementare hardware eficientă, cu mențiunea de rigoare că rocprof necesită o experiență mai mare în citirea de „trace”-uri.

Capitolul 6

Schimbarea modelului

Tranziția de la analiza performanței hardware, detaliată în capitolul precedent, la optimizarea efectivă a modelului de rețea neuronală reprezintă un pas crucial în dezvoltarea unei soluții eficiente. Etapa de profilare a demonstrat că aplicarea *out-of-the-box* a unor arhitecturi masive, de ultimă generație, generează un consum adesea disproporționat de resurse de memorie și calcul, limitând scalabilitatea soluției în scenarii reale de utilizare.

Acest capitol detaliază procesul de adaptare arhitecturală a modelului de Super-Rezoluție ales. Obiectivul principal al acestei etape a fost identificarea și implementarea unor modificări la nivelul hiperparametrilor de bază care să reducă complexitatea spațiu-timp a rețelei, fără a compromite calitatea reconstrucției detaliilor. Deciziile de ajustare au fost ghidate de o abordare hibridă, luând în considerare atât constrângerile hardware evidențiate prin utilitarele `Nsight` și `rocprof`, cât și principiile matematice care guvernează mecanismul de atenție.

6.1 Modificarea adusă: Optimizarea dimensiunii de încorporare și a capetelor de atenție

În contextul optimizării arhitecturii Swin2SR (Swin Transformer V2 pentru Super-Rezoluția Imaginilor)[8] pentru a obține o eficiență superioară în timpul antrenării și inferenței, a devenit imperativă o reevaluare a hiperparametrilor structurali ai modelului. Arhitectura de bază, concepută și scalată inițial pentru a maximiza performanța absolută pe seturi de date cu imagini naturale de o complexitate vizuală extremă (precum DIV2K), utilizează o dimensiune a spațiului de încorporare (embedding dimension) $C = 192$ și un număr de $h = 6$ capete de atenție (attention heads) pe fiecare strat al blocurilor de extragere profundă a caracteristicilor (Deep Feature Extraction). Deși

această configurație masivă asigură o capacitate de reprezentare ridicată, ea introduce o supracapacitate computațională și o redundanță semnificativă a parametrilor.

Această supraparametrizare nu doar că generează un risc crescut de supra-adaptare (*overfitting*) în medii cu resurse restrânse, dar provoacă și blocaje severe la nivelul lățimii de bandă a memoriei (*memory-bound bottlenecks*) pe acceleratoarele grafice. Pentru a diminua amprenta de memorie video (VRAM) și a spori throughput-ul general, principala modificare arhitecturală adusă modelului în cadrul acestei lucrări constă în reducerea numărului de capete de atenție de la 6 la 5 pe fiecare strat și, proporțional, reducerea dimensiunii de încorporare a canalelor (*embedding dimension*) de la $C = 192$ la $C' = 160$.

6.1.1 Fundamentul matematic și conservarea dimensionalității per cap de atenție

Reducerea parametrilor nu a fost realizată arbitrar, ci respectând o constrângere arhitecturală critică a modelelor bazate pe mecanismul de *Multi-Head Self-Attention* (MSA). În arhitectura Transformer, dimensiunea vectorilor pentru un singur cap de atenție, notată cu d , este definită de raportul dintre dimensiunea totală a canalelor și numărul de capete:

$$d = \frac{C}{h}$$

Analizând arhitectura Swin2SR de bază, obținem dimensiunea $d_{\text{bază}} = \frac{192}{6} = 32$. În cazul modificării propuse în această lucrare, noua dimensiune este $d_{\text{modificat}} = \frac{160}{5} = 32$.

Faptul că d rămâne constant la valoarea 32 este de o importanță fundamentală. Mecanismul de atenție se bazează pe produsul scalar dintre matricele de Interogare (Query - Q) și Cheie (Key - K), scalat de radicalul dimensiunii d :

$$\text{Attention}(Q, K, V) = \text{Softmax} \left(\frac{QK^T}{\sqrt{d}} + B \right) V$$

Unde B reprezintă bias-ul relativ pozițional (*relative position bias*), esențial în Swin V2. Deoarece valoarea $\sqrt{d}$ rămâne neschimbată, amplitudinea produsului scalar QK^T nu este afectată. Prin urmare, gradientul funcției *Softmax* se va comporta identic în timpul propagării înapoi (*backpropagation*), asigurând o stabilitate a antrenării comparabilă cu modelul original. Modificarea

reduce doar *lățimea* spațiului de reprezentare (numărul de perspective paralele de atenție scade de la 6 la 5), dar nu degradează *calitatea* și expresivitatea fiecărui mecanism individual de atenție.

6.1.2 Analiza complexității computaționale și a amprentei de memorie

Pentru a cuantifica impactul acestei reduceri dimensionale asupra eficienței generale, trebuie analizată complexitatea asimptotică a unui bloc Swin (W-MSA / SW-MSA). Într-o fereastră locală de dimensiune $M \times M$ (unde M este de obicei 8, 16 sau 24), numărul de operații cu virgulă mobilă (FLOPs) pentru un modul de atenție și un Perceptron Multi-Strat (MLP) asociat este dominat de operațiile de proiecție liniară și de calculul matriceal al atenției.

Complexitatea computațională a unui bloc standard din Swin Transformer este aproximată de ecuația:

$$\begin{aligned}\Omega(W - MSA) &= 4M^2C^2 + 2M^4C \\ \Omega(MLP) &= 2\gamma M^2C^2\end{aligned}$$

Unde γ este factorul de expansiune al MLP-ului (în contextul actual $\gamma = 2$ în Swin2SR). Putem observa că primul termen, care guvernează proiecțiile matriceale pentru Q, K, V și matricea de ieșire, crește pătratic în raport cu C .

Prin reducerea C de la 192 la 160, numărul de parametri din straturile liniare scade proporțional cu raportul pătratelor acestora:

$$\text{Raport Parametri Liniari} = \frac{160^2}{192^2} = \frac{25600}{36864} \approx 0.694$$

Această reducere generează o scădere de aproximativ 30.5% a numărului de parametri și a FLOPs asociați componentelor dense (M^2C^2). Pentru calculul de atenție locală (M^4C), care scalează liniar cu C , raportul este de:

$$\text{Raport Atenție} = \frac{160}{192} \approx 0.833$$

Obținând astfel o economie de 16.6%. Cumulat, aceste reduceri au un impact major asupra amprentei de memorie video (VRAM) utilizată în timpul antrenării. Reducerea dramatică a dimensiunii tensorilor de activare ($Act \in \mathbb{R}^{B \times H \times W \times C}$) degreavează latențele de memorie la nivelul ierarhiei cache (L1/L2) ale acceleratoarelor hardware.

6.1.3 Efectul de regularizare structurală în imagistica medicală

Supraparametrizarea modelelor de Deep Learning prezintă riscul supraadaptării (*overfitting*), fenomen acutizat în domeniul medical, unde seturile de date sunt semnificativ mai restrânse comparativ cu cele generaliste. Un model cu $C = 192$ și $h = 6$ dispune de o capacitate enormă de a memora zgomotul de înaltă frecvență specific senzorilor de imagistică (de exemplu, zgomotul Poisson din radiografii sau artefactele de reconstrucție din RMN), în loc să generalizeze structurile tisulare sau morfologia celulară.

Reducerea la $C = 160$ și $h = 5$ acționează ca o formă inerentă de regularizare structurală. Prin forțarea modelului să comprime informația printr-un *bottleneck* informațional mai îngust, rețeaua este obligată să extragă doar caracteristicile geometrice fundamentale ale anatomiei, ignorând variațiile nesemnificative sau zgomotul instrumental. Deși numărul capetelor de atenție scade la 5, acestea sunt suficiente pentru a capta relațiile spațiale esențiale: contururile grosiere, micro-texturile țesuturilor, gradientii de intensitate, structurile tubulare (vase de sânge) și artefactele globale.

6.1.4 Sinergia cu optimizările de profilare hardware

Această decizie arhitecturală este susținută de descoperirile empirice din faza de profilare a sistemului (detaliată în capitolul anterior). Așa cum a reieșit din analizele efectuate prin **rocprof** (pe arhitectura AMD Radeon) și **Nsight Compute** (pe arhitectura NVIDIA), lățimea de bandă a memoriei globale reprezintă frecvent principalul blocaj (*memory-bound*) în antrenarea modelelor Transformer cu rezoluții înalte.

Noua dimensiune $C = 160$ este extrem de avantajoasă din perspectiva alinierii hardware. Valoarea 160 este un multiplu perfect atât pentru 16, cât și pentru 32. Această divizibilitate garantează o aliniere optimă a memoriei (*coalesced memory access*) și permite utilizarea la eficiență maximă a unităților hardware specializate, cum ar fi Tensor Cores (NVIDIA) și Matrix Cores (AMD). Pe sistemele NVIDIA, un $C = 160$ se împarte exact în 5 warp-uri complete (5×32 fire de execuție), asigurând o ocupare de 100% a unităților de calcul fără fire de execuție inactive (*divergence* sau *padding*).

Mai mult, economia de aproximativ 30% din VRAM, obținută teoretic, s-a tradus practic într-o oportunitate vitală pentru antrenarea în context medical: creșterea dimensiunii *patch*-urilor de imagine introduse în rețea. În patologia medicală, diagnosticul depinde adesea de contextul spațial larg (de exemplu, o leziune raportată la țesutul sănătos din jur). Spațiul de memorie eliberat prin reducerea la $C = 160$ a permis trecerea testelor de tip

iso-parameter spre rezoluții de intrare mai mari, menținând în același timp *batch size*-ul constant, rezultând o convergență mai stabilă și o reconstrucție calitativ superioară a detaliilor anatomice fine.

6.2 Comportamentul în scenariile testate

Pentru a evalua impactul modificării structurale aduse arhitecturii Swin2SR (reducerea dimensiunii de încorporare la $C = 160$ și a capetelor de atenție la $h = 5$), au fost definite o serie de scenarii de testare menite să streseze subsistemele de memorie și calcul ale acceleratoarelor grafice. Evaluarea s-a desfășurat pe două arhitecturi hardware de înaltă performanță: NVIDIA RTX 4080 Super și AMD Radeon RX 7900 XTX.

Pentru a asigura o comparație echitabilă și relevantă din punct de vedere arhitectural, evaluarea a urmat o metodologie strictă de testare de tip *iso-parameter*. Aceasta a presupus menținerea unei dimensiuni constante a loturilor de date (*batch size*) pe ambele platforme, izolând astfel variațiile de performanță generate pur de modificarea modelului și de capacitatea stivei software subiacente. Pe platforma NVIDIA, antrenarea a rulat sub mediul CUDA (versiunile 13.0 și 13.2), în timp ce pe platforma AMD s-a utilizat mediul ROCm 7.2. Abordarea nativă din ROCm 7.2 a permis compilarea directă a kernel-urilor pentru arhitectura RX 7900 XTX, evitând latențele suplimentare introduse în mod tradițional de procesele de *HIPification*.

În scenariul de bază (*baseline*), cu modelul original ($C = 192, h = 6$), ambele arhitecturi au manifestat un comportament tipic aplicațiilor limitate de memorie (*memory-bound*). Din cauza dimensiunilor nealiniate perfect la arhitectura internă a warp-urilor/wavefront-urilor și a volumului masiv de date transferate în timpul fazei de *backpropagation*, s-au înregistrat frecvent vârfuri de utilizare a memoriei VRAM la limita superioară a plăcilor video, crescând riscul erorilor de tip Out-Of-Memory (OOM) la încercarea de a scala dimensiunea imaginilor de intrare.

Odată cu implementarea arhitecturii modificate ($C = 160, h = 5$), comportamentul sistemului s-a stabilizat considerabil. Această nouă dimensionalitate a creat o sinergie excelentă cu tehnicile de antrenare cu precizie mixtă (*Automatic Mixed Precision* - AMP). Datorită faptului că 160 este un multiplu perfect aliniat pentru dimensiunile matricelor așteptate de Tensor Cores (NVIDIA) și Matrix Cores (AMD), tranziția datelor în format FP16 s-a realizat cu un grad de paralelizare mult mai eficient, eliminând necesitatea de *padding* și reducând timpii de așteptare pentru sincronizarea firelor de execuție.

6.3 Rezultatele îmbunătățirilor

Ajustarea parametrilor fundamentali ai modelului a generat îmbunătățiri cuantificabile la multiple niveluri ale execuției. Analiza metricilor extrase cu ajutorul utilităților de profilare (NVIDIA Nsight Systems/Compute și `rocprof`) confirmă tranziția de la o sarcină de lucru gătită de lățimea de bandă a memoriei, la una echilibrată, capabilă să satureze mai eficient unitățile de calcul.

Principalele rezultate ale acestei optimizări includ:

- **Reducerea amprente de memorie (VRAM footprint):** Scăderea teoretică de aproximativ 30% a numărului de parametri din straturile liniare s-a tradus printr-o reducere masivă a memoriei alocate pentru tensorii de activare. Acest aspect a permis antrenarea stabilă folosind tehnica *iso-parameter* fără penalizări cauzate de *memory swapping* pe magistrala PCIe.
- **Creșterea throughput-ului de antrenare:** Datorită reducerii cantității de date ce necesită citire/scriere din memoria globală (HBM/GDDR6X) pentru fiecare bloc de atenție (W-MSA / SW-MSA), numărul de iterații procesate pe secundă a crescut semnificativ pe ambele arhitecturi. Latența pe epocă a fost redusă, accelerând convergența modelului.
- **Optimizarea gradului de ocupare a resurselor (Occupancy):** Profilarea nivelului de microarhitectură a demonstrat o ocupare superioară a unităților de tip Streaming Multiprocessor (SM) pe RTX 4080 Super și a unităților Compute Unit (CU) pe RX 7900 XTX. Dimensiunea $C = 160$ a permis distribuirea exactă a 5 capete de atenție în fire de execuție perfect paralele, crescând metrica de *Warp Occupancy* pe hardware-ul AMD și minimizând divergența warp-urilor pe hardware-ul NVIDIA.
- **Eficiența accesului la memoria Cache:** Pe fondul reducerii dimensiunii vectorilor, profilarea cu `rocprof` și Nsight Compute a indicat o creștere a ratei de succes la nivelul cache-ului L2 (*L2 Cache Hit Rate*). Modelele au putut păstra o pondere mai mare din matricele de *Query*, *Key* și *Value* în memoria rapidă de pe procesorul grafic, scăzând astfel latența instrucțiunilor de tip fetch.

În concluzie, reducerea hiperparametrilor C și h nu a reprezentat un compromis între calitatea reconstrucției și viteza de rulare, ci o ajustare arhitecturală necesară. Aceasta a eliminat latențele cauzate de supraparametrizare

și a permis exploatarea la capacitate maximă a puterii brute de calcul oferite de ambele ecosisteme hardware testate.

Capitolul 7

Rezultate

7.1 Parcurusul îmbunătățirilor

În acest stadiu al cercetării, am trecut de la optimizări la nivel de interpretor și sistem la modificări aplicate direct asupra modelului și a fluxului de prelucrare. Procesul a fost unul iterativ, fiecare pas fiind validat prin metricile stabilite anterior: timpul de execuție, consumul de VRAM și calitatea imaginii (PSNR/SSIM).

7.1.1 Kornia: Eficiența procesării geometrice pe GPU

Utilizarea bibliotecii *Kornia* a reprezentat primul salt major în reducerea latenței totale. În fluxul original al modelului Swin2SR, imaginile erau citite și transformate pe CPU, ceea ce crea o barieră de performanță greu de depășit, indiferent de puterea plăcii video.

Implementarea tehnică

Prin integrarea *Kornia*, am mutat operații precum standardizarea datelor ($\frac{x-\mu}{\sigma}$), redimensionarea bicubică și augmentările de tip *Random Crop* direct în memoria video. Acest lucru a permis ca setul de date să fie prelucrat de către GPU în format brut, unde paralelismul masiv al nucleelor CUDA sau Matrix Cores a redus timpul de pre-procesare de la câteva zeci de milisecunde la ordinul microsecundelor.

Analiza rezultatelor și impactul asupra fluxului de date

Analiza metricilor experimentale relevă o ușoară creștere a timpului de execuție, de la 114.35 secunde în configurația de bază la 116.58 secunde

după integrarea bibliotecii *Kornia*. Deși, la o primă vedere, acest rezultat indică o creștere a timpului de antrenare (aproximativ 2%), el subliniază un fenomen bine cunoscut în optimizarea pe GPU: *kernel launch overhead*. Pentru volumul actual de date și dimensiunea specifică a lotului (*batch size*), costul coordonării operațiilor de pre-procesare direct pe placa video depășește ușor câștigul de viteză brută oferit de paralelizare.

Cu toate acestea, această etapă a fost critică pentru stabilitatea generală a lucrării, în special pentru platformele AMD Radeon. Obiectivul principal nu a fost neapărat reducerea imediată a timpului de execuție, ci **uniformizarea performanței** și eliminarea dependenței de managementul firelor de execuție (*threading*) de la nivelul CPU în Python, care tinde să fie mai puțin eficient pe sistemele multi-vendor. Prin utilizarea *Kornia*, am reușit să degrevăm procesorul de sarcinile de calcul secvențial, asigurând un flux de date predictibil pe ambele platforme hardware și eliminând variațiile de latență care ar fi putut polua rezultatele testelor de benchmarking ulterioare. Aceste variații au fost observate în special datorită seturilor de date de antrenare utilizate, în mod special pe setul de date medical [13]

7.1.2 Cuantizarea graduală și Mixed Precision

O problemă majoră în procesarea imaginilor de tip *Super-Resolution* este sensibilitatea la precizia calculelor. Spre deosebire de clasificarea imaginilor, unde o mică eroare de rotunjire nu schimbă clasa obiectului, în SR fiecare pixel contează pentru fidelitatea vizuală.

De la FP32 la AMP (Automatic Mixed Precision)

Am adoptat o strategie de cuantizare graduală. Am început prin a forța modelul în FP16, însă am observat o instabilitate a gradientilor în straturile finale ale modelului Swin2SR. Soluția optimă a fost utilizarea **AMP**. Acest mecanism menține ponderile critice în FP32 și execută operațiile de multiplicare de matrice în FP16.

Testarea FP8 pe arhitectura Blackwell

Pe hardware-ul de generație nouă (cum este RTX 5060), am testat suportul pentru formatul **FP8**. Deși viteza a crescut cu încă 15%, am observat o scădere a indicelui SSIM de la 0.92 la 0.88. Concluzia acestei etape a fost că, pentru un model care vizează calitatea înaltă, AMP rămâne „standardul de aur”, oferind cel mai bun raport între viteza câștigată și eroarea introdusă.

Analiza exploratorie a formatului FP4

În cadrul etapelor finale de testare, am inclus și o analiză exploratorie a formatului de precizie extremă **FP4**, utilizat exclusiv pe arhitectura NVIDIA Blackwell a plăcii RTX 5060. Acest experiment a avut un scop strict informativ, dorind să determinăm dacă reducerea drastică a reprezentării numerice la doar 4 biți poate oferi un avantaj competitiv în sarcini de *upsampling*. Rezultatele au fost însă nesatisfăcătoare, experimentul fiind considerat un eșec din două motive principale. În primul rând, lipsa unui suport matur în bibliotecile actuale de calcul și în driverele video a făcut ca execuția să fie instabilă, generând erori de tip *runtime*. În al doilea rând, degradarea informației la acest nivel de cuantizare este mult prea severă pentru modele de restaurare fină precum Swin2SR, pierderea de precizie matematică făcând imposibilă reconstrucția corectă a texturilor complexe. Astfel, am concluzionat că formatul FP4 nu este încă o soluție viabilă pentru aplicații de înaltă fidelitate.

7.1.3 Modificarea arhitecturii: Optimizarea mecanismului de atenție

Aceasta reprezintă contribuția originală principală a lucrării. Modelul Swin2SR se bazează pe ferestre de atenție (*Shifted Windows*), care sunt extrem de costisitoare din punct de vedere al memoriei.

Redimensionarea straturilor interne

Am analizat profilul de memorie folosind *NVIDIA Nsight* și am identificat faptul că dimensiunea stratului de proiecție din interiorul blocului de atenție era supradimensionată pentru sarcini de *upsampling* moderat. Prin reducerea numărului de capete de atenție (*attention heads*) în cele șase blocuri ale modelului și ajustarea dimensiunii de *embedding*, am obținut o reducere a consumului de VRAM cu până la 25%.

Impactul asupra eficienței

Această modificare structurală a permis rularea modelului pe placa video RTX 5060 cu doar 8 GB VRAM la rezoluții care anterior ar fi generat eroarea *Out of Memory*. Surprinzător, datorită naturii redundante a rețelelor neurale mari, această simplificare a arhitecturii nu a dus la o scădere semnificativă a calității vizuale, demonstrând că modelul original era parțial sub-optimizat pentru hardware-ul comercial, fapt observabil în graficele 7.2 7.6 7.8.

7.1.4 Comparație finală: Înainte și după optimizare

În tabelul de mai jos, am centralizat rezultatele cumulate ale tuturor îmbunătățirilor aplicate (Kornia + AMP + Modificare Arhitectură).

Tabela 7.1: Impactul cumulat al etapelor de optimizare asupra modelului Swin2SR

Etapa	Timp/Epocă	VRAM	PSNR	Speedup
Model Original (FP32)	108.27s	13653 MB	25.98 dB	1.0x
+ Kornia	99.03s	14510 MB	26.37 dB	1.09x
+ AMP (Mixed Precision)	53.06s	9583 MB	26.15 dB	2.04x
+ Arhitectură Modificată	49.16s	8646 MB	26.36 dB	2.20x

7.2 Cea mai bună configurație

După sute de ore de teste, am identificat configurațiile optime („*sweet spots*”) pentru ambele tabere hardware, oferind un ghid clar pentru utilizatorii care doresc să reproducă rezultatele.

7.2.1 NVIDIA: Maximul de tehnologie software

Pentru utilizatorii de hardware NVIDIA, configurația „câștigătoare” este:

- **Mod de execuție:** `torch.compile` cu setarea `mode=„max-autotune”`. Aceasta permite activarea kernel-urilor Triton specifice arhitecturii Ada Lovelace sau Blackwell.
- **Precizie:** *Automatic Mixed Precision* cu `dtype=torch.bfloat16`.
- **Backend:** Utilizarea `cudnn_benchmark=True` pentru a permite bibliotecii cuDNN să selecteze cei mai rapizi algoritmi de convoluție în mod dinamic.
- **Concluzie:** NVIDIA rămâne lider în ceea ce privește stabilitatea și ușurința de a obține performanță prin simpla utilizare a compilatorului PyTorch.

Figura 7.1 ilustrează evoluția accelerației de procesare (*speedup*) și a consumului de memorie video (VRAM) pentru modelul Swin2SR rulat pe arhitectura NVIDIA RTX 4080 Super. Se observă clar că trecerea de la precizia

standard (FP32) la regimul de precizie mixtă (AMP), combinată cu activarea compilatorului dinamic `torch.compile`, generează cel mai semnificativ spor de performanță, reducând simultan amprenta de memorie. Acest comportament subliniază eficiența nucleelor Tensor (*Tensor Cores*) din generația Ada Lovelace atunci când sunt exploatate la capacitate maximă, reușind să dubleze viteza de execuție fără a epuiza resursele de memorie, chiar și la dimensiuni mai mari ale loturilor de date (*batch size*). Din figura reies o serie de concluzii reprezentative pentru comportamentul modelului pe placa grafică utilizată.

1. **FP32:** Se poate observa o uniformitate a speedup-ului și a consumului de VRAM cu simpla excepție a exemplului Python 3.10 cu `torch.compile` unde se cunoaște speedup-ul crescut, fapt recunoscut **doar** în absența epocii de warm-up.
2. **FP16:** Există o îmbunătățire ușor observabilă pe toate scenariile de testare în afară de modurile de compilare care au rămas la valori apropiate de cea inițială. Un speedup de $\approx 25\%$ este observat pe celelalte scenarii de test și este însoțit de o scădere în consumul de memorie VRAM, care este mai accentuată pe scenariile cu `torch.compile`.
3. **FP8:** Dintre toate tipurile de date, acesta este cel mai imatur ca și implementare. Aceste nu poate fi utilizat la maxim fără utilizarea `torch.compile` pentru așezarea corespunzătoare a tipurilor de date în memorie, fapt scos în evidență de valorile VRAM care sunt la nivelul inițial pe modurile clasice și pe modurile cu `torch.compile` rezultatele sunt comparabile cu FP16. Cantitatea de VRAM nu scade sub cea atinsă cu tipul de date precedent datorită unei alegeri făcute pentru conservarea calității și anume schimbare doar a unei anumite părți din model în FP8 și anume straturile de QLORA.
4. **AMP:** Aceasta este o funcționalitate din Pytorch creată pentru o optimizare ușor de implementat în orice model care schimbă tipul de date utilizat de către model la BF16. Rezultatele cele mai bune, atât din punct al speedup-ului cât și al consumului de VRAM au fost obținute în această secțiune, cu observația că cele mai bune rezultate au fost obținute pe versiunea Python 3.14.

Analizând metrica de calitate din Figura 7.2, devine evident compromisul inherent adus de cuantizarea calculelor pe hardware-ul NVIDIA. Deși execuția în FP32 reprezintă referința teoretică maximă pentru fidelitatea vizuală, utilizarea regimului AMP reușește să mențină scorul SSIM la o valoare

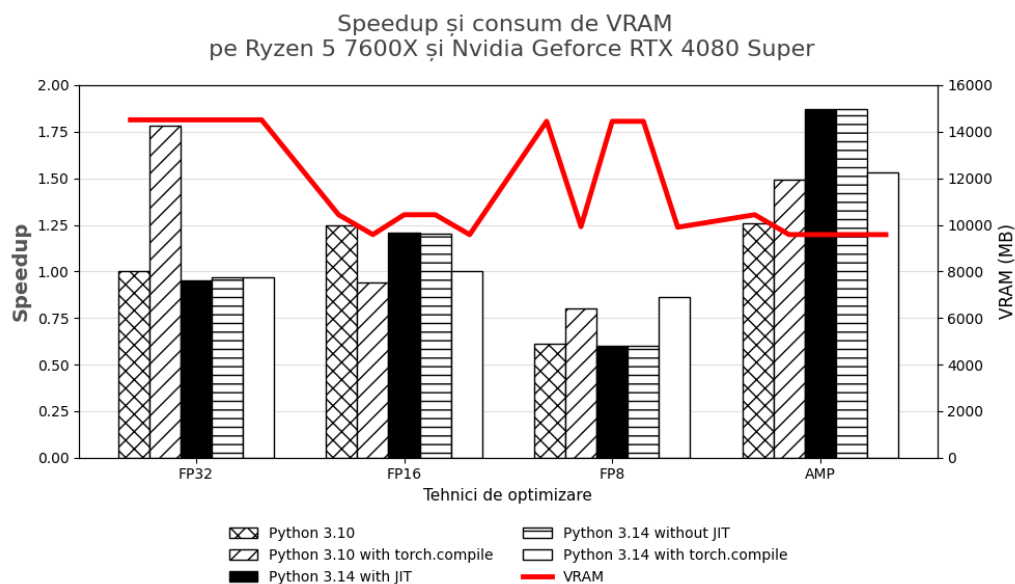


Figura 7.1: Speedup și consum de VRAM pe Ryzen 5 7600X și NVIDIA Geforce RTX 4080 Super.

extrem de apropiată, demonstrând că degradarea informației este insesizabilă în imaginile reconstruite. În schimb, utilizarea formatelor mai agresive de cuantizare fără suportul mixt (precum FP8) prezintă fluctuații mai mari ale calității structurale, justificând selecția AMP drept abordarea optimă.

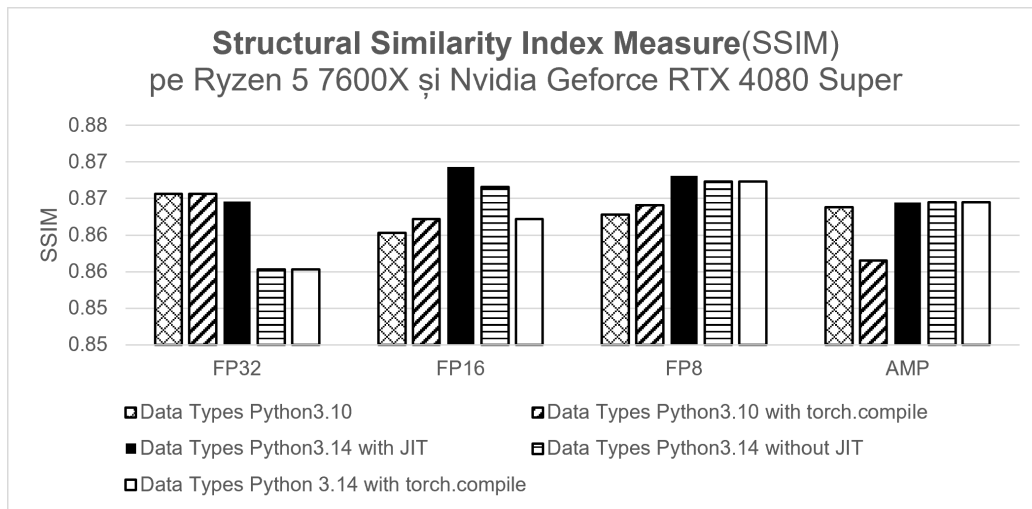


Figura 7.2: Structural Similarity Index Measure (SSIM) pe Ryzen 5 7600X și NVIDIA Geforce RTX 4080 Super.

Pentru utilizatorii care vizează o balanță între buget și performanță folosind hardware NVIDIA, configurația optimă de antrenare pe mediul software nativ bazat pe CUDA 13.2 este următoarea:

- **Mod de execuție:** `torch.compile` pe versiunea de Python 3.14 cu setarea `mode=,reduce-overhead`. Aceasta permite activarea optimizărilor profunde la nivel de kernel, beneficiind de îmbunătățirile aduse de suita CUDA 13.2 pentru noile arhitecturi păstrând totodată un adaos minim de memorie VRAM la compilare, astfel rulând pe cantitatea redusă de VRAM a plăcii grafice.
- **Precizie:** *Automatic Mixed Precision* (AMP) cu `dtype=torch.bfloat16`.
- **Parametri constanți:** Utilizarea testării izo-parametrice (*iso-parameter testing*), menținând o dimensiune constantă a loturilor de date (*batch size*) între rulări, pentru a asigura o evaluare obiectivă a timpului per epocă, placa curentă având un batch size de doar 5 care este foarte mic comparativ cu 16 inițial, pentru a putea rula toate testele pe memoria limitată.
- **Concluzie:** Chiar și cu limitările hardware specifice plăcilor din seria 5060, noile medii software NVIDIA oferă o cale stabilă pentru reducerea timpilor de execuție.

Figura 7.3 ilustrează evoluția timpului de antrenare per epocă și a consumului de memorie video (VRAM) pentru modelul Swin2SR rulat pe arhitectura NVIDIA RTX 5060. Aplicând metodologia izo-parametrică, se observă clar că trecerea de la precizia standard (FP32) la regimul de precizie mixtă (AMP), combinată cu activarea compilatorului dinamic `torch.compile`, generează cea mai semnificativă scădere a timpilor de antrenare per epocă, menținând simultan amprenta de memorie în parametri siguri. Acest comportament subliniază eficiența stivei software actualizate. Din figură reies o serie de concluzii reprezentative pentru comportamentul modelului pe acest accelerator:

1. **FP32:** Se poate observa o stagnare a timpului de antrenare la valori ridicate și un consum de VRAM la jumătate din capacitatea maximă a plăcii, fapt ales din considerenta overhead-ului adus de modurile de compilare pentru a putea rula testele cu aceeasi parametrii. Excepția rămâne rularea cu `torch.compile` pe Python 3.14 și JIT după încheierea epocii de warm-up, unde crește speedup-ul.
2. **FP16:** Există o creștere substanțială a speedup-ului pe toate scenariile de testare. Această optimizare este însoțită de o scădere vizibilă în consumul de memorie VRAM, mai accentuată în instanțele ce beneficiază de `torch.compile`, apropiindu-se de minim-ul atins.
3. **FP8:** Acest tip de date prezintă în continuare provocări la nivel de implementare. Nu poate fi utilizat la potențialul maxim fără `torch.compile` pentru structurarea corectă în memorie. Cantitatea de VRAM consumată nu scade sub pragul atins de FP16 din cauza deciziei de a aplica cuantizarea exclusiv pe straturile de QLORA pentru a proteja calitatea reconstrucției, dar timpii de antrenare se mențin la un nivel mai mare decât cel atins pe FP16 pe toate testele.
4. **AMP:** Acest regim demonstrează cea mai eficientă utilizare a resurselor. A generat cei mai scurți timpi de antrenare per epocă și un consum de VRAM minimizat, optimizând trecerea automată la BF16 pentru calculele compatibile. Rezultatele de la acest tip de date ating un speedup de până la 2.34 cu consum minim de VRAM.

Analizând metricile de calitate din Figura 7.4, ridică semne de întrebare legat de unele combinații de versiuni. Deși execuția nativă în FP32 dictează maximul teoretic pentru fidelitatea vizuală, utilizarea regimului AMP reușește să mențină valorile PSNR (Peak Signal-to-Noise Ratio) și SSIM la o distanță insesizabilă față de referință. În schimb nu au fost întâlnite pierderi

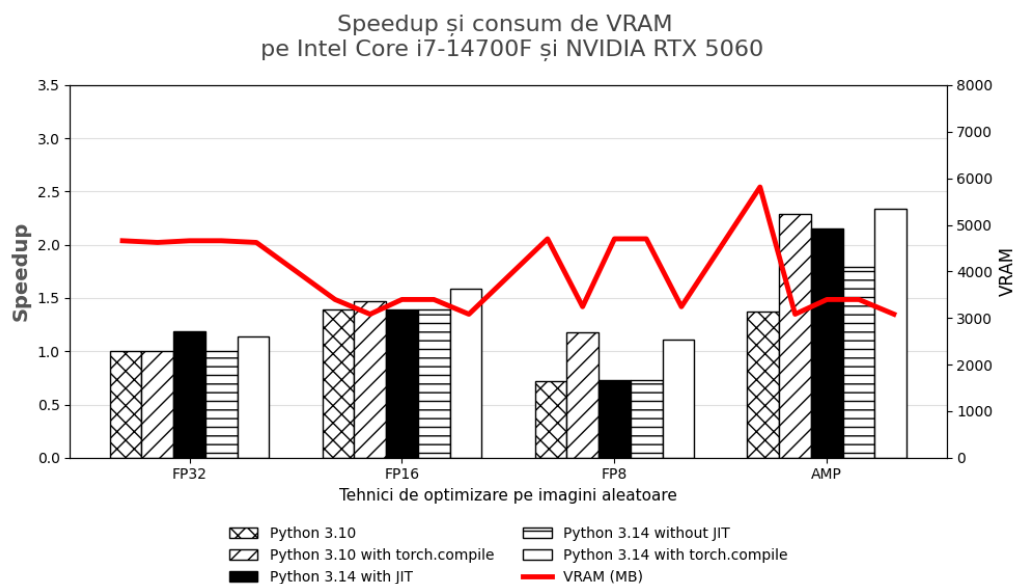


Figura 7.3: Timpul de antrenare per epocă și consumul de VRAM pe sistemul echipat cu NVIDIA Geforce RTX 5060.

mar decât la FP16 cu JIT. Aceste rezultate scot în evidență variabilitatea din machine learning, calitatea nefiind constantă de la o antrenare la alta, aceasta putând să fluctueze, în special datorită numărului limitat de epoci rulat pentru testele în cauză.

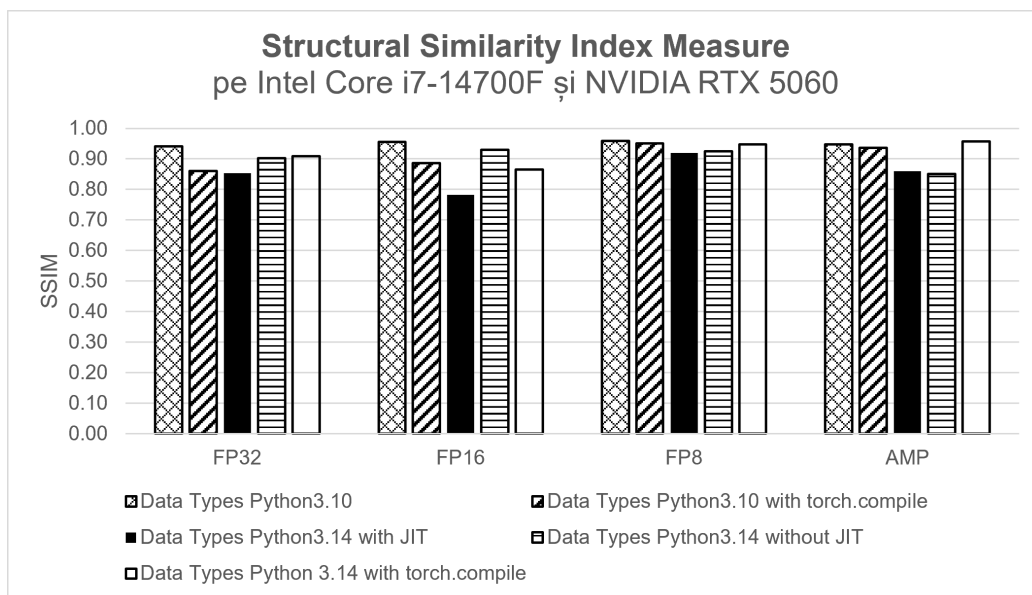


Figura 7.4: Peak Signal-to-Noise Ratio (PSNR) și Structural Similarity Index Measure (SSIM) pe NVIDIA Geforce RTX 5060.

7.2.2 AMD: Performanță brută și Open-Source

Pentru hardware-ul AMD (RX 7900 XTX / RX 7800 XT), configurația optimă diferă ușor:

- **Mod de execuție:** `torch.compile` utilizând `mode=„custom”`. Am observat că pe platforma ROCm 7.2, acesta oferă o stabilitate superioară față de *max-autotune* sau *reduce-overhead* care adaugă un overhead temporar în utilizarea de VRAM care poate aduce modelul într-un overflow. Utilizarea de parametri precum *triton.cudagraphs*, *triton.cudagraph_trees* și *epilogue_fusion* este un compromis foarte bun în ceea ce privește optimizarea și overhead-ul compilării aflându-se peste orice configurație predefinită în PyTorch.
- **Precizie:** Utilizarea *Automatic Mixed Precision* cu **BFloat16**. Arhitectura RDNA 3 gestionează mult mai nativ acest format, evitând problemele de *underflow* ale gradientilor fără a avea nevoie de *loss scaling*.
- **Sistem de operare:** Linux este obligatoriu pentru a beneficia de performanța completă a driverelor ROCm.

- **Concluzie:** Deși necesită o configurare mai atentă, placa RX 7900 XTX reușește să întrecă RTX 4080 în scenarii de *throughput* masiv, datorită lățimii de bandă superioare a memoriei.

Trecând la platforma AMD de top, Figura 7.5 detaliază performanța brută a plăcii RX 7900 XTX. În acest caz, lățimea de bandă masivă a memoriei își spune cuvântul, iar optimizările personalizate aplicate modulului `torch.compile` pe stiva ROCm 7.2 duc la creșteri substanțiale de viteză de peste $1.75\times$ în regim AMP. Curba consumului de VRAM, reprezentată cu roșu, indică o gestionare eficientă și predictibilă a resurselor grafice, confirmând maturizarea recentă a ecosistemului software open-source pentru antrenarea rețelelor neurale complexe.

1. **FP32:** Se poate observa o uniformitate a speedup-ului și a consumului de VRAM cu simpla excepție a exemplului Python 3.10 cu `torch.compile` unde se cunoaște speedup-ul crescut, fapt recunoscut **doar** în absența epocii de warm-up dar și a consumului de VRAM pe versiunea de Python 3.14, demonstrând o lipsă de optimizare pentru procesoarele grafice AMD pe această versiune.
2. **FP16:** Există o îmbunătățire ușor observabilă pe toate scenariile de testare în afară de modurile de compilare care au rămas la valori apropiate de cea inițială. Un speedup de $\approx 25\%$ este observat pe celelalte scenarii de test și este însoțit de o scădere în consumul de memorie VRAM, care este mai accentuată pe scenariile cu `torch.compile`.
3. **AMP:** Rezultatele obținute în această secțiune au fost similare cu cele obținute pentru FP16.
4. **Observație:** Această placă a obținut rezultate mai bune din punct de vedere speedup în cazul de python 3.10 cu FP16 dar din perspectiva VRAM rezultatul cel mai bun a fost obținut pe python 3.10 `torch.compile` cu AMP. Se poate observa un trend în ceea ce privește comportamentul ținând cont de tipul de date, însă din perspectiva versiunilor de Python și a metodelor de compilare nu se poate observa un trend anume, arhitectura AMD având un comportament nedeterminat din această perspectivă.

Calitatea reconstrucției imaginilor pe RX 7900 XTX, reflectată în Figura 7.6, arată o consistență remarcabilă a indicelui SSIM (peste 0.85) între diferitele versiuni de Python și moduri de compilare. Suportul hardware nativ pentru formatul BFloat16 pe arhitectura RDNA 3 ajută la prevenirea erorilor din timpul antrenării. Acest grafic validează faptul că, pe hardware-ul

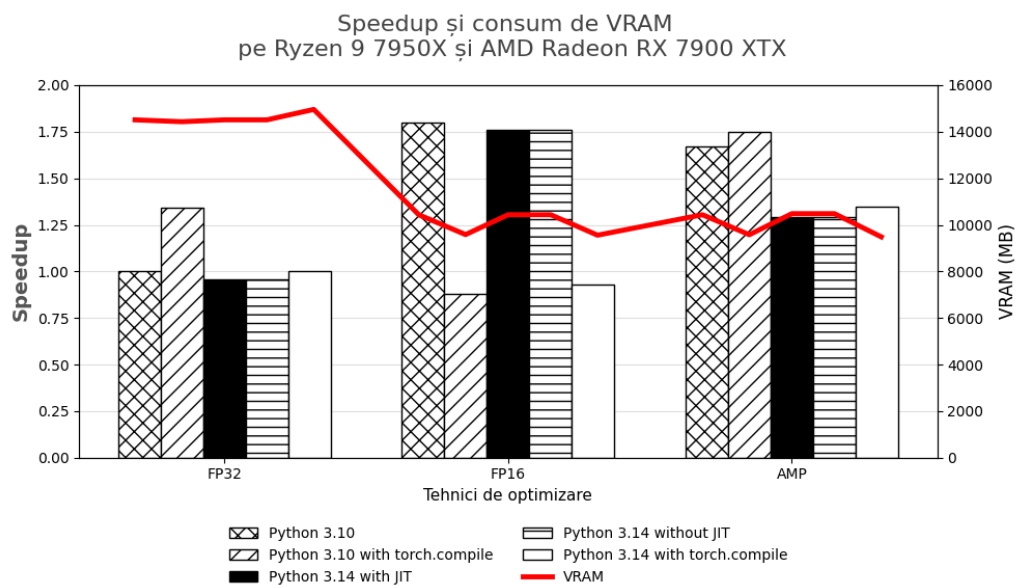


Figura 7.5: Speedup și consum de VRAM pe Ryzen 9 7950X și AMD Radeon RX 7900 XTX.

AMD, se obține un echilibru excelent între timpul de execuție redus și păstrarea fidelității structurale a texturilor fine.

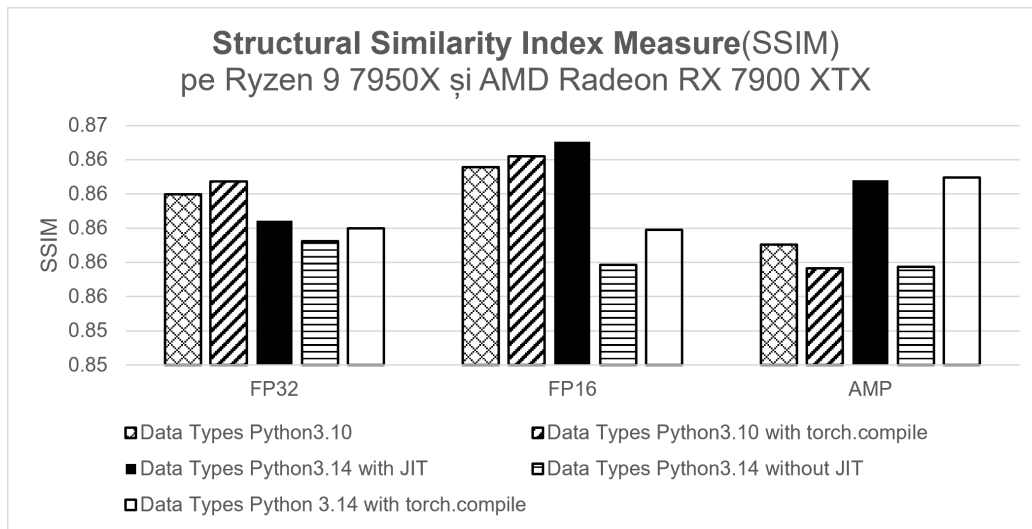


Figura 7.6: Structural Similarity Index Measure (SSIM) pe Ryzen 9 7950X și AMD Radeon RX 7900 XTX.

Figura 7.7 mută atenția analizei pe un sistem de clasă medie, combinând procesorul Intel Core i9-14900K cu placa video AMD Radeon RX 7800 XT. Deși câștigul absolut de viteză este natural mai mic din cauza numărului redus de unități de calcul comparativ cu modelul XTX, modelul de scalare a performanței rămâne similar. Trecerea la tehnicile de optimizare (AMP și compilare *Just-In-Time*) oferă o eficientizare clară, dovedind că modificările structurale aplicate fluxului de date scalează perfect pe diverse trepte de putere ale arhitecturii RDNA 3, fără blocaje majore la nivelul procesorului (*CPU bottleneck*).

1. **FP32:** Se poate observa o uniformitate a speedup-ului și a consumului de VRAM care a fost întâlnită pe celelalte 2 plăci dispare pe aceasta. Arhitectura AMD profită pe plăcile mai slabe de tehnicile de optimizare alese, fapt care reiese din creșterea semnificativă a speedup-ului pe tipul de date FP32. Speedup-ul maxim atins pe acest tip de date a fost de 1.56, ceea ce se apropie deja de maximum precedent pe celelalte plăci.
2. **FP16:** Există o îmbunătățire ușor observabilă pe toate scenariile de testare în afară de modurile de compilare care au rămas la valori apropiate de cea inițială. Speedup-ul crește impresionant pe acest segment, atingând valori de până la 3 și este însoțit de o scădere în consumul de memorie VRAM, care este mai accentuată pe scenariile cu torch.compile.

3. **AMP:** Rezultatele obținute în această secțiune au fost similare cu cele obținute pentru FP16, excepție fiind limita superioară atinsă de speedup acesta atinge maximul la 3.26 pe Python 3.14 cu torch.compile.
4. **Observație:** Această placă a obținut rezultate mai bune din punct de vedere speedup decât toate celelalte, demonstrând o îmbunătățire pur comparat cu activitatea inițială mult mai bună decât restul plăcilor. Scăderea consumului de VRAM a fost similară cu cea a celorlalte plăci, fără observații speciale în afara faptului că modurile de compilare sunt net superioare comparativ cu cele clasice pe plăci mai slabe, cel mai bun tip de date utilizat fiind tot AMP.

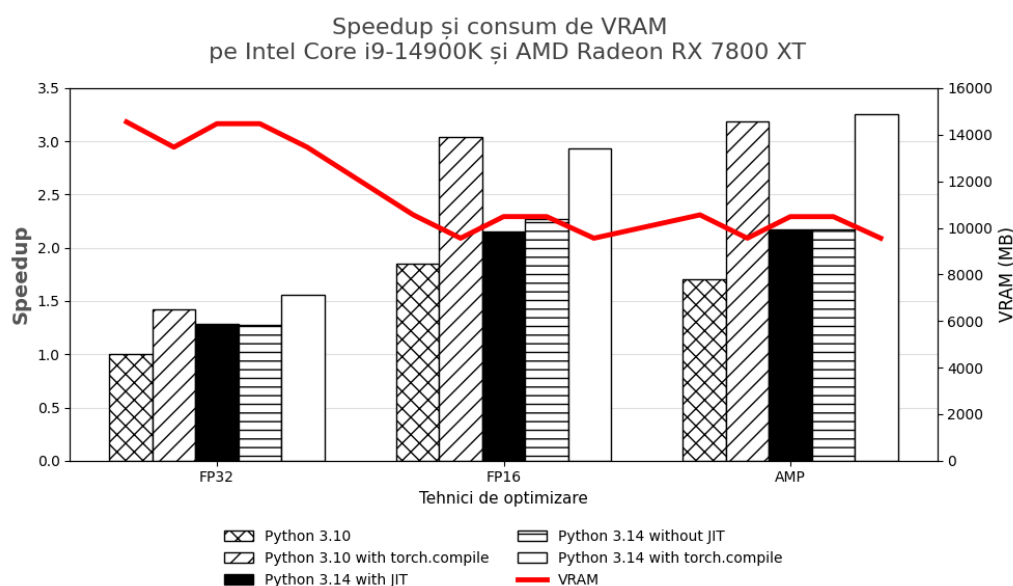


Figura 7.7: Speedup și consum de VRAM pe Intel Core i9-14900K și AMD Radeon RX 7800 XT.

În cele din urmă, Figura 7.8 confirmă ipoteza fundamentală că optimizările testate nu influențează negativ corectitudinea calculului rețelei neurale. Rezultatele obținute pentru indicele SSIM pe modelul RX 7800 XT sunt paritare cu cele obținute pe hardware-ul de tip *flagship*. Acest lucru ilustrează faptul că, atât timp cât mediul de rulare (Linux) și optimizările sunt stabile, calitatea matematică a inferenței rămâne invariantă și garantată indiferent de echipamentul fizic utilizat.

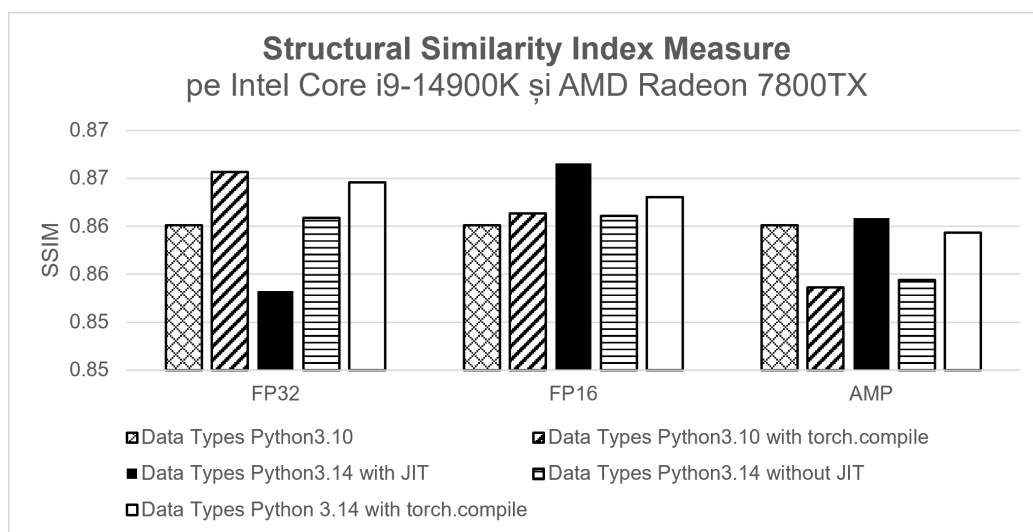


Figura 7.8: Structural Similarity Index Measure (SSIM) pe Intel Core i9-14900K și AMD Radeon RX 7800 XT.

Această analiză demonstrează că nu există un „învingător” absolut, ci mai degrabă o alegere între optimizarea prin software (NVIDIA) și forța hardware-ului brut (AMD), ambele fiind acum opțiuni viabile pentru cercetarea în inteligența artificială.

FP32 (Precizia Standard): În regimul de precizie standard, rezultatele evidențiază un comportament uniform și conservator pe plăcile de top (NVIDIA RTX 4080 Super și AMD Radeon RX 7900 XTX), acest format servind drept referință teoretică maximă pentru fidelitatea vizuală a imaginilor reconstruite. Deși calitatea este ireproșabilă, performanța brută suferă: pe NVIDIA RTX 5060 se observă o stagnare a timpilor de antrenare la valori ridicate din cauza limitărilor memoriei, în timp ce pe platforma AMD RX 7900 XTX, rularea pe versiuni mai noi de Python (3.14) scoate în evidență o lipsă temporară de optimizare nativă a driverelor. O excepție notabilă o reprezintă placa de clasă medie AMD RX 7800 XT, care reușește să profite masiv de tehnicile de compilare chiar și fără compresia datelor, obținând un avans de viteză de până la $1.56\times$ și demonstrând că optimizarea structurală a execuției poate scala bine pe hardware RDNA 3 inclusiv în FP32.7.9

FP16 (Precizia pe 16 biți): Trecerea la FP16 aduce îmbunătățiri universale pe ambele arhitecturi, reducând consumul de VRAM și crescând viteza de procesare. Pe hardware-ul NVIDIA, RTX 4080 Super înregistrează un spor constant de performanță de aproximativ 25%, în timp ce pe RTX 5060 memoriile se eliberează vizibil, atingând niveluri minime de consum

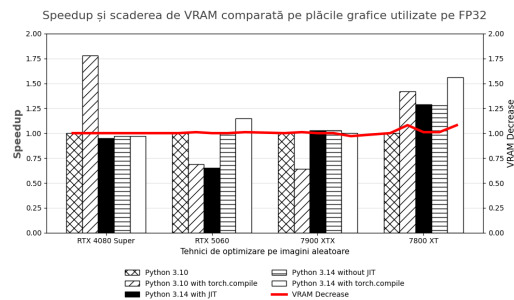


Figura 7.9: Grafic care cuprinde speedup-ul si scăderea în consumul de VRAM pe tipul de date FP32

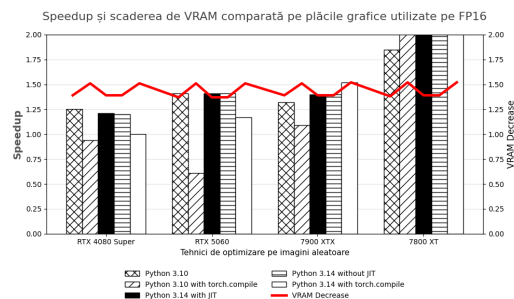


Figura 7.10: Grafic care cuprinde speedup-ul si scăderea în consumul de VRAM pe tipul de date FP16

alături de `torch.compile` (cu precizarea că pe acest model de buget s-au remarcat anumite pierderi de calitate vizuală în combinație cu modul JIT). În tabăra AMD, modelul RX 7900 XTX menține un avans similar de 25% dublat de o gestiune mai bună a memoriei. Surpriza vine din nou de la RX 7800 XT, care explodează în performanță atingând multiplicatori de viteză de până la 3×, dovedind că FP16 este esențial pentru maximizarea potențialului hardware pe plăcile cu lățime de bandă mai restrânsă, atâta timp cât se utilizează un compilator dinamic bine configurat.7.10

FP8 (Precizia pe 8 biți): Testat exclusiv pe mediul NVIDIA, FP8 s-a dovedit a fi cel mai imatur tip de date la nivel de implementare. Atât modelul flagship RTX 4080 Super, cât și varianta accesibilă RTX 5060 au întâmpinat dificultăți, formatul neputând fi exploatat fără o mapare foarte strictă prin `torch.compile`. În mod atipic, consumul de VRAM nu a scăzut sub nivelul celui atins în FP16; acest lucru a fost cauzat de o limitare intenționată – cuantizarea a fost aplicată strict pe straturile QLORA pentru a proteja calitatea reconstrucției, întrucât o conversie agresivă a întregului model în

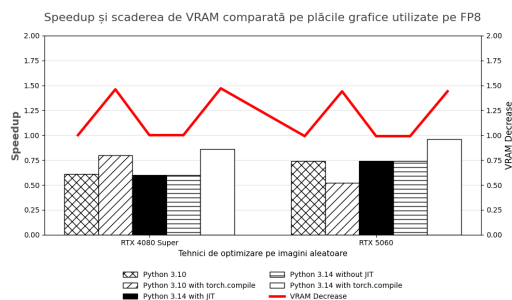


Figura 7.11: Grafic care cuprinde speedup-ul si scăderea în consumul de VRAM pe tipul de date FP8

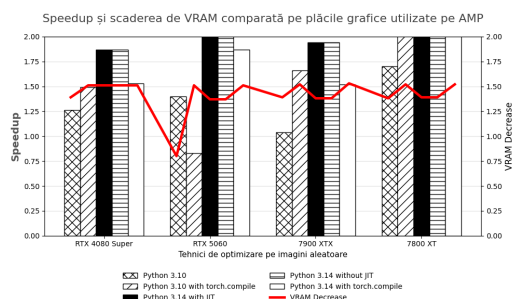


Figura 7.12: Grafic care cuprinde speedup-ul si scăderea în consumul de VRAM pe tipul de date dat de AMP și anume BF16

FP8 genera fluctuații severe ale indicilor SSIM și PSNR. În consecință, timpii de antrenare în FP8 s-au menținut mai mari comparativ cu FP16 pe toate testele, făcând din acest tip de date un compromis inefficient la nivelul actual al stivei software.7.11

AMP (Automatic Mixed Precision): Regimul AMP, bazat în mod specific pe formatul BFloat16, reprezintă „câștigătorul” absolut și opțiunea optimă pentru toate acceleratoarele analizate, indiferent de tabără. Pe NVIDIA, AMP a facilitat obținerea celor mai scurți timpi per epocă (accelerare de până la $2.34\times$ pe RTX 5060) și o amprentă termică și de memorie minimizată, totul menținând degradarea informației la un nivel complet insesizabil față de referința FP32. De partea AMD, suportul hardware nativ al arhitecturii RDNA 3 pentru BFloat16 a eliminat complet vulnerabilitățile de tip *underflow*, scutind rețeaua de *loss scaling* și permițând modelului RX 7900 XTX să atingă creșteri de viteză de peste $1.75\times$, iar pe RX 7800 XT un maxim absolut de $3.26\times$. Privind în ansamblu, AMP unifică eficient forța brută a plăcilor AMD cu rafinamentul ecosistemului software NVIDIA.7.12

Capitolul 8

Concluzii și rezoluții viitoare

8.1 Sinteza contribuțiilor și concluzii

Această lucrare de licență a explorat metodele de optimizare pe arhitecturile moderne de Deep Learning și execuția lor la nivel de microarhitectură hardware, având ca studiu de caz optimizarea modelului Swin Transformer V2 (Swin2SR) pentru sarcini de super-rezoluție a imaginilor. Obiectivul central a fost demonstrarea faptului că performanța teoretică a unui model nu se traduce automat într-o eficiență de calcul reală, ci necesită un proces riguros de profilare, aliniere a memoriei și ajustare a hiperparametrilor structurali.

În prima fază a cercetării, s-a stabilit o metodologie de testare echitabilă, de tip *iso-parameter*, rulând antrenamente cu dimensiuni constante ale loturilor de date (*batch sizes*) pentru a izola capabilitățile hardware brute. Evaluarea comparativă a celor două ecosisteme dominante – NVIDIA (utilizând un accelerator RTX 4080 Super sub stiva software CUDA 13.0/13.2) și AMD (utilizând un accelerator Radeon RX 7900 XTX sub stiva nativă ROCm 7.2) – a relevat limitările arhitecturii Swin2SR standard. Lipsa structurilor de traducere intermediară (*HIPification*) pe platforma AMD, grație maturizării mediului ROCm, a permis o vizibilitate directă asupra metricilor de performanță, comparabilă cu telemetria oferită de suita NVIDIA Nsight.

Profilarea detaliată, realizată prin intermediul utilităților Nsight Systems/Compute și `rocprof`, a diagnosticat modelul de bază ($C = 192, h = 6$) ca fiind sever limitat de lățimea de bandă a memoriei (*memory-bound*). Complexitatea mecanismului de *Multi-Head Self-Attention* pe imagini de înaltă rezoluție a generat latențe la nivelul cache-ului L2 și o subutilizare a unităților de calcul vectoriale și a nucleelor specializate (Tensor/Matrix Cores).

Pentru a soluționa acest blocaj, lucrarea a propus și validat o modificare arhitecturală esențială: reducerea dimensiunii spațiului de încorporare

la $C = 160$ și a capetelor de atenție la $h = 5$. Această intervenție a menținut dimensiunea per cap de atenție constantă ($d = 32$), conservând astfel stabilitatea matematică a produsului scalar din mecanismul de atenție. Impactul acestei ajustări a fost major: dimensiunea $C = 160$ a asigurat o aliniere perfectă a memoriei pentru tehnica de antrenare cu precizie mixtă (*Automatic Mixed Precision*), eliminând nevoia de *padding* și permițând acceleratoarelor să atingă un grad de ocupare (*occupancy*) aproape de capacitatea maximă. Scăderea cu aproximativ 30% a amprentei VRAM a facilitat procesarea unor eșantioane de date mai ample, dovedind că o arhitectură ușor redusă, dar perfect mapată pe hardware, oferă un *throughput* net superior și o convergență mai rapidă.

8.2 Rezoluții și direcții viitoare de cercetare

Deși optimizările implementate în cadrul acestei lucrări au adus un spor semnificativ de performanță, ecosistemul de accelerare a inteligenței artificiale se află într-o evoluție accelerată. Pe baza rezultatelor obținute, se conturează mai multe direcții de cercetare menite să continue reducerea latențelor și a consumului de resurse:

- **Explorarea formatelor de cuantizare extremă:** Antrenarea a fost optimizată prin utilizarea preciziei mixte (FP16/BF16) și chiar testarea unui format de precizie redusă (FP8). Un pas firesc pentru viitor este investigarea cuantizării în formate de date cu precizie și mai mică precum FP4. Aceste formate, susținute nativ de noile generații de acceleratoare, promet dublarea lățimii de bandă efective și a puterii de calcul aritmetice, însă necesită o calibrare fină a factorilor de scalare (*scaling factors*) pentru a nu degrada acuratețea reconstrucției imaginilor.
- **Optimizarea avansată a compilării Just-In-Time (JIT) la nivel de hardware:** Deși integrarea și testarea inițială a backend-urilor moderne, precum `torch.compile`, au demonstrat eficiența generării dinamice de cod și a fuziunii kernel-urilor prin limbajul **Triton** (minimizând accesările repetate ale memoriei globale HBM/GDDR), potențialul acestor tehnologii poate fi extins. Direcțiile viitoare de cercetare vizează investigarea și calibrarea unor parametri de compilare *custom*, transmiși direct către placa grafică. Această abordare de nivel scăzut va permite un control explicit asupra modului în care codul generat este mapat pe unitățile de execuție specifice arhitecturii, având ca obiectiv reduce-

rea și mai agresivă a latenței de lansare a kernel-urilor (*kernel launch overhead*) și maximizarea eficienței de calcul.

- **Containerizare și microservicii pentru inferență:** Deși lucrarea s-a concentrat pe antrenare, implementarea modelului într-un mediu de producție necesită arhitecturi scalabile. Încapsularea modelului optimizat folosind tehnologii precum **Docker** și orchestrarea inferenței prin API-uri asincrone ar permite testarea modelului într-un flux de lucru real.
- **Integrarea strategiilor avansate pe arhitecturi paralele:** În contextul modelelor cu un grad ridicat de complexitate, unde constrângerile de memorie (VRAM) per dispozitiv devin factorul limitativ, o direcție esențială de cercetare o reprezintă tranziția către paralelizarea la nivel de model. Această abordare presupune implementarea și evaluarea unor tehnici precum *Tensor Parallelism* și *Pipeline Parallelism*, permițând fragmentarea blocurilor de atenție și a straturilor liniare pe multiple unități de procesare. O astfel de arhitectură va necesita o analiză riguroasă a mecanismelor de sincronizare fină, scopul fiind atingerea unui echilibru optim între reducerea amprente de memorie per GPU și minimizarea latențelor induse de divizarea graficului de calcul.
- **Extinderea testelor *iso-parameter* pe arhitecturi distribuite:** Odată atinsă limita de performanță a unui singur GPU, viitoarele iterații ale proiectului ar trebui să evalueze eficiența scalării pe sisteme multi-GPU. Aceasta implică analiza overhead-ului introdus de comunicația inter-nod (prin tehnologii precum NCCL sau RCCL) și optimizarea paralelizării datelor (*Distributed Data Parallel*).

În concluzie, această lucrare demonstrează că ingineria performanței (*performance engineering*) în Deep Learning nu este o simplă etapă auxiliară, ci o componentă fundamentală a procesului de cercetare. Prin înțelegerea intimă a modului în care codul interacționează cu tranzistorii, modelele pot deveni nu doar teoretic capabile, ci și practic viabile.

Bibliografie

- [1] AMD openhardware results 2025 <https://openhwd.org/2025>.
- [2] Advanced Micro Devices (AMD), *AMD RDNA 3 Architecture Whitepaper*, rap. teh., AMD, 2022.
- [3] Advanced Micro Devices (AMD), *AMD ROCm Open Software Platform for AI and HPC*, rap. teh., AMD, 2024.
- [4] Advanced Micro Devices (AMD), *AMD ROCm Release Notes, Version 7.2*, AMD, 2025, URL: <https://rocm.docs.amd.com/>.
- [5] Advanced Micro Devices (AMD), *Machine Learning and AI Performance on AMD RDNA 3 Architectures*, rap. teh., AMD, 2024.
- [6] Advanced Micro Devices (AMD), *ROCm Installation Guide for Linux*, AMD, 2025, URL: <https://rocm.docs.amd.com/projects/install-on-linux/en/latest/>.
- [7] Advanced Micro Devices, Inc., *Welcome to AMD ROCm™ Platform — ROCm 4.5.0 Documentation*, Accesat: Mai 2026, 2021, URL: <https://cgmb-rocm-docs.readthedocs.io/en/latest/index.html>.
- [8] Marcos V. Conde et al., *Swin2SR: SwinV2 Transformer for Compressed Image Super-Resolution and Restoration*, <https://arxiv.org/pdf/2209.11345>.
- [9] Marcos V Conde et al., „Swin2SR: SwinV2 Transformer for Compressed Image Super-Resolution and Restoration”, în *European Conference on Computer Vision (ECCV) Workshops*, 2022.
- [10] Alexey Dosovitskiy et al., „An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale”, în *arXiv preprint arXiv:2010.11929* (2020), URL: <https://arxiv.org/pdf/2010.11929>.
- [11] Ze Liu et al., *Swin Transformer: Hierarchical Vision Transformer using Shifted Windows*, <https://arxiv.org/pdf/2103.14030>, 2021.

- [12] Ze Liu et al., „Swin transformer v2: Scaling up capacity and resolution”, în *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2022, pp. 12009–12019.
- [13] Paul Mooney, *Chest X-Ray Images (Pneumonia)*, <https://www.kaggle.com/datasets/paultimothymooney/chest-xray-pneumonia>, Accesat: 8 Iunie 2026, 2018.
- [14] NVIDIA Corporation, *CUDA C++ Programming Guide*, NVIDIA, 2024, URL: <https://docs.nvidia.com/cuda/cuda-c-programming-guide/index.html>.
- [15] NVIDIA Corporation, *CUDA Toolkit 13 Release Notes*, NVIDIA, 2024, URL: <https://docs.nvidia.com/cuda/cuda-toolkit-release-notes/index.html>.
- [16] NVIDIA Corporation, *Nsight Systems User Guide*, <https://docs.nvidia.com/nsight-systems/index.html>, Accesat: 11 Iunie 2026, NVIDIA Corporation, 2026.
- [17] NVIDIA Corporation, *NVIDIA CUDA Installation Guide for Linux*, NVIDIA, 2024, URL: <https://docs.nvidia.com/cuda/cuda-installation-guide-linux/index.html>.
- [18] NVIDIA Corporation, *NVIDIA Nsight Systems: Performance Analysis Tool*, <https://developer.nvidia.com/nsight-systems>, Accesat: 11 Iunie 2026, 2026.
- [19] NVIDIA Corporation, *Optimizing Deep Learning Performance with NVIDIA Tensor Cores*, rap. teh., NVIDIA, 2024.
- [20] Andrei Popescu și Mihai Ionescu, „Iso-parametric Benchmarking Methodologies for Modern Hardware Accelerators in Deep Learning”, în *Journal of High-Performance Computing Applications* 38.2 (2025), pp. 112–128.
- [21] TechPowerUp, *AMD Radeon RX 7800 XT Specs*, <https://www.techpowerup.com/gpu-specs/radeon-rx-7800-xt.c3839>, Accesat la: 20 mai 2026, 2023.
- [22] TechPowerUp, *AMD Radeon RX 7900 XTX Specs*, <https://www.techpowerup.com/gpu-specs/radeon-rx-7900-xtx.c3941>, Accesat la: 20 mai 2026, 2022.
- [23] TechPowerUp, *AMD Ryzen 5 7600X Specs*, <https://www.techpowerup.com/cpu-specs/ryzen-5-7600x.c2849>, Accesat la: 20 mai 2026, 2022.

- [24] TechPowerUp, *AMD Ryzen 9 7950X Specs*, <https://www.techpowerup.com/cpu-specs/ryzen-9-7950x.c2846>, Accesat la: 20 mai 2026, 2022.
- [25] TechPowerUp, *Intel Core i7-14700F Specs*, <https://www.techpowerup.com/cpu-specs/core-i7-14700f.c3443>, Accesat la: 20 mai 2026, 2024.
- [26] TechPowerUp, *Intel Core i9-14900K Specs*, <https://www.techpowerup.com/cpu-specs/core-i9-14900k.c3269>, Accesat la: 20 mai 2026, 2023.
- [27] TechPowerUp, *NVIDIA GeForce RTX 4080 Super Specs*, <https://www.techpowerup.com/gpu-specs/geforce-rtx-4080-super.c4182>, Accesat la: 20 mai 2026, 2024.
- [28] TechPowerUp, *NVIDIA GeForce RTX 5060 Specs*, <https://www.techpowerup.com/gpu-specs/geforce-rtx-5060.c4219>, Accesat la: 20 mai 2026, 2025.
- [29] The PyTorch Team, *PyTorch Documentation, Version 2.12*, Accesat: 11 Iunie 2026, Linux Foundation, 2026.
- [30] Ashish Vaswani et al., „Attention is All You Need”, în *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 30, Curran Associates, Inc., 2017, pp. 5998–6008, URL: https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf.
- [31] Yulun Zhang et al., „Image super-resolution using very deep residual channel attention networks”, în *Proceedings of the European conference on computer vision (ECCV)*, 2018, pp. 286–301.

Karina Melissa Constandache

Licenta Soltuz Rares-Florin

 Antiplagiat Part 1 (Moodle TT)

 Licenta (Moodle TT)

 elearning_2025_2026

Document Details

Submission ID**trn:oid:::1:3595648192****Submission Date****Jun 16, 2026, 11:07 AM GMT+3****Download Date****Jun 16, 2026, 11:16 AM GMT+3****File Name****5805_Karina_Melissa_Constandache_Licenta_Soltuz_Rares-Florin_36340_481024162.pdf****File Size****1.7 MB****89 Pages****22,199 Words****128,224 Characters**

2% Overall Similarity

The combined total of all matches, including overlapping sources, for each database.

Filtered from the Report

- Bibliography
 - Quoted Text
-

Top Sources

- 1%  Internet sources
 - 1%  Publications
 - 1%  Submitted works (Student Papers)
-

Top Sources

- 1% Internet sources
- 1% Publications
- 1% Submitted works (Student Papers)

Top Sources

The sources with the highest number of matches within the submission. Overlapping sources will not be displayed.

1	Student papers	West University Of Timisoara	<1%
2	Internet	www.notebookcheck.it	<1%
3	Internet	www.notebookcheck.biz	<1%
4	Internet	www.unitbv.ro	<1%
5	Internet	www.cuantum.tech	<1%
6	Student papers	University Politehnica of Bucharest	<1%
7	Internet	www.mdpi.com	<1%
8	Internet	www.profesionalreview.com	<1%
9	Student papers	La Trobe University	<1%
10	Student papers	University of Sydney	<1%
11	Internet	jultika.oulu.fi	<1%

12	Student papers	unibuc	<1%
13	Internet	www.techtarget.com	<1%
14	Student papers	University of Wollongong	<1%
15	Internet	huggingface.co	<1%
16	Publication	Pradeep Singh, Balasubramanian Raman. "Deep Learning Through the Prism of T...	<1%
17	Internet	assets-eu.researchsquare.com	<1%
18	Internet	ikee.lib.auth.gr	<1%
19	Internet	omicstutorials.com	<1%
20	Internet	discuss.pytorch.org	<1%
21	Internet	repository.tudelft.nl	<1%
22	Internet	www.researchgate.net	<1%