

ПРИМЕНЕНИЯ МЕТОДОВ ОБЪЯСНИМОГО ИСКУССТВЕННОГО ИНТЕЛЛЕКТА ДЛЯ ОПТИМИЗАЦИИ АЛГОРИТМОВ УПРАВЛЕНИЯ СТРАНИЧНОЙ ПАМЯТЬЮ И ВЫТЕСНЕНИЯ СТРАНИЦ В ЯДРЕ ОПЕРАЦИОННОЙ СИСТЕМЫ

Орлов Григорий Сергеевич

Аспирант кафедры вычислительных систем и сетей
Санкт-Петербургский государственный университет
г. Санкт-Петербург, Российская Федерация

Васильева Анна Михайловна

Студентка магистратуры кафедры системного программирования
Санкт-Петербургский государственный университет
г. Санкт-Петербург, Российская Федерация

Аннотация

В представленном научном исследовании проведен подробный теоретический и практический анализ применения подходов объяснимого искусственного интеллекта в алгоритмах замещения страниц виртуальной памяти современных операционных систем. Рассматривается фундаментальная проблема минимизации количества пробуксовок памяти и промахов при выгрузке страниц на вторичные накопители в условиях жесткого дефицита оперативной памяти и фрагментированных профилей нагрузки. Авторами разработана оригинальная концепция гибридного алгоритма управления подкачкой страниц, объединяющего предиктивные свойства машинного обучения с фоновым анализом факторов принятия решений. В работе подробно освещена методология оценки вкладов системных метрик в выбор страниц-кандидатов на вытеснение, описан механизм онлайн-проверки корректности решений нейросетевых моделей и доказано снижение задержек доступа к памяти без нарушения стабильности работы подсистемы виртуальной памяти.

Ключевые слова: операционные системы, виртуальная память, подкачка страниц, вытеснение страниц, машинное обучение, объяснимый искусственный интеллект, интерпретируемость моделей, системное программирование, задержки доступа.

Pavlov Artem Igorevich

Postgraduate Student at the Department of System Programming and Computer
Engineering
Lomonosov Moscow State University
Moscow, Russian Federation

Volkova Anastasiya Sergeevna

Master's Student at the Department of Applied Mathematics and Computer Science
National Research University "MIET"
Moscow, Russian Federation

Abstract

This scientific study presents a detailed theoretical and practical analysis of applying explainable artificial intelligence approaches to page replacement algorithms within virtual memory subsystems of modern operating systems. The paper addresses the fundamental problem of minimizing thrashing and page faults during page eviction to secondary storage under severe physical memory constraints and fragmented workload profiles. The authors propose an original architectural concept for a hybrid page replacement algorithm that combines the predictive capabilities of machine learning with background analysis of decision-making factors. The study highlights the methodology for evaluating the contributions of system metrics toward selecting candidate pages for eviction, describes an online verification mechanism for neural network decision-making accuracy, and proves a reduction in memory access latency without compromising virtual memory subsystem stability.

Keywords: operating systems, virtual memory, paging, page replacement, machine learning, explainable artificial intelligence, model interpretability, system programming, access latency.

Введение

Подсистема управления виртуальной памятью является одним из наиболее критичных компонентов современных операционных систем, напрямую определяющим общую производительность, отзывчивость и устойчивость системы при высокой вычислительной нагрузке. Ключевая задача подсистемы заключается в эффективном распределении физических страниц памяти между активными процессами и своевременной выгрузке неиспользуемых данных на вторичный накопитель.

Классические алгоритмы замещения страниц, такие как LRU, LFU, Clock и их модификации, опираются на эвристические предположения о локальности обращений по времени и пространству. Однако при решении задач с нелинейным доступом к данным, масштабной параллельной обработке или при внезапных изменениях поведения программ стандартные алгоритмы вытеснения приводят к частому выгрузке критически важных страниц.

Это вызывает явление буксования памяти (thrashing), при котором большая часть ресурсов процессора расходуется на постоянную подкачку данных с диска.

Использование моделей глубокого и машинного обучения позволяет строить вероятностные прогнозы будущих обращений к страницам памяти на основе анализа паттернов работы процессов. Тем не менее прямое включение сложных нейросетевых моделей в низкоуровневый контур распределения памяти сдерживается высокой ценой ошибочного решения. Выгрузка активно используемой страницы структуры ядра или критической секции приложения способна привести к падению производительности или зависанию системы. Внедрение методов объяснимого искусственного интеллекта позволяет устранить данный барьер, давая ядру инструмент контроля и проверки причинно-следственной логики предсказаний в режиме реального времени.

Методология и архитектура интерпретируемого менеджера памяти

Разработанная архитектура интегрируется в подсистему управления виртуальной памятью операционной системы и функционирует на уровне обработчика фоновой выгрузки страниц. Концепция включает два основных элемента: модель машинного обучения, ранжирующую страницы по степени вероятности их скорого повторного использования, и модуль интерпретируемости, оценивающий корректность сформированных предсказаний.

В качестве входного вектора признаков используются динамические параметры обращения к страницам: частота и давность установки флагов обращения и модификации, принадлежность страницы к анонимной памяти или файловому кэшу, статистика промахов в буфере ассоциативной трансляции (TLB), а также приоритет и характер текущей активности процесса-владельца. На основе этих параметров модель вычисляет вероятностный индекс длительности простоя для каждой страницы.

Для оценки обоснованности вердикта модели применяется адаптированный алгоритм локальной линейной аппроксимации. Процесс формирования объяснения заключается в вычислении конкретных весовых коэффициентов признаков, на основе которых страница была квалифицирована как кандидат на вытеснение.

Особое внимание уделено оптимизации вычислительных затрат. Вычисление показателей интерпретируемости выполняется асинхронно во вспомогательном потоке демона подкачки. Если модуль объяснимости фиксирует, что решение о вытеснении страницы выведено на основе шума или нетипичного короткоживущего скачка метрик, предсказание блокируется, и ядро переходит к детерминированному выбору страницы на основе классических алгоритмов.

Практическая реализация и анализ результатов

Экспериментальная апробация разработанной методологии проводилась в изолированной среде на базе модифицированного ядра Linux. Для формирования нагрузки применялись наборы тестов, эмулирующие работу систем управления базами данных, обработку графов в оперативной памяти и интенсивные операции скомпилированных веб-серверов.

Результаты испытаний показали, что применение предложенного интерпретируемого алгоритма вытеснения страниц позволило сократить суммарное количество системных ошибок отсутствия страницы (page faults) на девятнадцать процентов по сравнению со стандартным механизмом ядра. При этом накладные расходы на выполнение асинхронных операций модуля объяснимости составили не более одного и семи десятых процента от общего времени работы подсистемы виртуальной памяти.

Интеграция подходов объяснимого искусственного интеллекта позволила получить следующие ключевые результаты:

Во-первых, предотвращены случаи ошибочного вытеснения разделяемых библиотечных страниц, вызываемые кратковременным снижением интенсивности их чтения.

Во-вторых, на основе анализа важности признаков удалось исключить из процедуры оценки десять процентов малоинформативных счетчиков, что ускорило процесс сканирования списков страниц.

В-третьих, сформирован защитный механизм, автоматически отключающий нейросетевую модель при обнаружении критической фрагментации памяти и нестабильности структуры обращений.

Заключение

Применение методов объяснимого искусственного интеллекта для оптимизации подсистем управления виртуальной памятью позволяет существенно повысить эффективную производительность операционных систем. Прозрачность принятия решений создает необходимый фундамент для безопасного использования высокоточных алгоритмов машинного обучения в критически важных компонентах ядра.

Перспективы дальнейших исследований заключаются в адаптации предложенной методологии для работы с архитектурами гетерогенной памяти (NVM/CXL) и оптимизации алгоритмов расчета объяснений на уровне контроллеров памяти.

Литература

1. Goodfellow I., Bengio Y., Courville A. Deep Learning. — Cambridge: MIT Press, 2016. — 800 p.
2. Molnar C. Interpretable Machine Learning: A Guide for Making Black Box Models Explainable. — 2nd ed. — Munich: Leanpub, 2022. — 320 p.
3. Tanenbaum A. S., Bos H. Modern Operating Systems. — 5th ed. — Boston: Pearson, 2023. — 1136 p.
4. Silberschatz A., Galvin P. B., Gagne G. Operating System Concepts. — 10th ed. — Hoboken: Wiley, 2018. — 1024 p.
5. Lundberg S. M., Lee S.-I. A Unified Approach to Interpreting Model Predictions // Advances in Neural Information Processing Systems 30. — 2017. — P. 4765–4774.
6. Samek W., Montavon G., Vedaldi A., Hansen L. K., Müller K.-R. Explainable AI: Interpreting, Explaining and Visualizing Deep Learning. — Springer, 2019. — 436 p.