

ПОСТРОЕНИЯ ИНТЕРПРЕТИРУЕМЫХ МОДЕЛЕЙ ГЛУБОКОГО ОБУЧЕНИЯ ДЛЯ ПРЕДИКТИВНОЙ ПРЕФЕТЧИНГ-ОПТИМИЗАЦИИ В ПОДСИСТЕМАХ ВВОДА-ВЫВОДА ОПЕРАЦИОННЫХ СИСТЕМ

Морозов Артем Викторович

Аспирант кафедры вычислительной техники и системного программирования
Московский физико-технический институт
(национальный исследовательский университет)
г. Москва, Российская Федерация

Алексеева Дарья Сергеевна

Студентка магистратуры кафедры информатики и вычислительной техники
Московский физико-технический институт
(национальный исследовательский университет)
г. Москва, Российская Федерация

Аннотация

В представленном фундаментальном научном труде осуществляется всесторонний теоретический и практический анализ интеграции интеллектуальных подсистем предиктивного подкачивания и кэширования данных в подсистемы ввода-вывода современных операционных систем. Рассматривается актуальная проблема сглаживания высокой задержки при обращении к энергонезависимым носителям информации и распределенным хранилищам данных в условиях нестационарного и труднопредсказуемого профиля запросов. Авторами разработана оригинальная методологическая концепция гибридного препетчера на базе нейросетевых моделей глубокого обучения, способного анализировать многомерные временные ряды системных вызовов и прогнозировать будущие обращения к блочным устройствам. Ключевой акцент исследования сделан на преодолении проблемы «черного ящика» с помощью адаптивных методов объяснимого искусственного интеллекта. В работе детально описаны механизмы локальной интерпретации решений модели в режиме реального времени, проанализировано влияние накладных расходов на вычисление вкладов признаков и доказана высокая эффективность предотвращения ложного подкачивания данных с сохранением стабильности работы дисковой подсистемы.

Ключевые слова: операционные системы, подсистема ввода-вывода, префетчинг, кэширование, машинное обучение, объяснимый искусственный интеллект, интерпретируемость моделей, блочные устройства, системное программирование.

Ivanov Sergey Alekseevich

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

Petrova Elena Mikhailovna

Master's Student at the Department of Computer Engineering
National Research University "MIET"
Moscow, Russian Federation

Abstract

This fundamental scientific paper carries out a comprehensive theoretical and practical analysis of integrating intelligent predictive prefetching and data caching subsystems into the input/output subsystems of modern operating systems. The study addresses the pressing issue of mitigating high latency when accessing non-volatile storage media and distributed data stores under non-stationary and hard-to-predict request profiles. The authors have developed an original methodological concept for a hybrid prefetcher based on deep learning neural network models capable of analyzing multidimensional system call time series and predicting future requests to block devices. A key emphasis of the research is placed on overcoming the "black box" problem using adaptive Explainable Artificial Intelligence (XAI) methods. The paper describes in detail real-time local decision interpretation mechanisms for the model, analyzes the impact of computational overhead from feature contribution calculations, and proves high efficiency in preventing false data prefetching while maintaining disk subsystem stability.

Keywords: operating systems, input/output subsystem, prefetching, caching, machine learning, explainable artificial intelligence, model interpretability, block devices, system programming.

Введение

Развитие современной архитектуры вычислительных систем сопровождается постоянным углублением разрыва между производительностью центральных процессоров и задержками доступа к подсистемам хранения данных. Несмотря на широкое внедрение высокоскоростных твердотельных накопителей и энергонезависимой памяти, операции ввода-вывода остаются одним из главных узких мест, определяющих общую отзывчивость и пропускную способность операционных систем при выполнении высоконагруженных задач.

Традиционные механизмы упреждающего чтения данных, интегрированные в подсистемы ввода-вывода современных ядер, преимущественно базируются на простых детерминированных эвристиках. Данные алгоритмы выявляют последовательные обращения к смежным блокам диска и инициируют их упреждающую загрузку в оперативную память.

Однако при работе со сложными структурами данных, базами данных или при одновременном выполнении множества параллельных процессов паттерны доступа становятся нелинейными и фрагментированными. В таких условиях классические эвристики либо теряют эффективность, либо приводят к эффекту избыточного подкачивания, бессмысленно расходуя пропускную способность шины и замусоривая страницы системного кэша.

Внедрение моделей глубокого обучения, таких как рекуррентные нейронные сети и архитектуры на основе механизмов внимания, позволяет строить высокоточные прогнозы будущих обращений на основе длинных цепочек исторических запросов. Тем не менее широкое применение нейросетей в подсистемах ядра сдерживается высокой сложностью их внутренней структуры и непредсказуемостью решений. Ошибочный или необоснованный префетчинг может вызвать вытеснение из памяти действительно необходимых страниц или привести к перегрузке контроллера накопителя. Использование подходов объяснимого искусственного интеллекта позволяет решить эту проблему, обеспечивая постоянный контроль за обоснованностью формируемых моделью прогнозов.

Методология и архитектура интерпретируемого префетчера

Архитектура предлагаемого интеллектуального модуля упреждающего чтения встраивается непосредственно в слой абстракции блочных устройств ядра операционной системы. Система включает в себя два ключевых компонента: предсказательный нейросетевой модуль, функционирующий в режиме реального времени, и фоновый модуль анализа интерпретируемости, осуществляющий непрерывную верификацию вырабатываемых гипотез.

В качестве входного признакового пространства используются пространственные и временные характеристики потока запросов ввода-вывода. К ним относятся идентификаторы вызвавших процессы потоков, логические номера запрашиваемых блоков, объемы считываемых данных, временные интервалы между смежными вызовами, а также текущее состояние заполненности буферного кэша и очередей команд дискового контроллера. На основе этих параметров модель прогнозирует вероятностное распределение обращений к блокам на несколько шагов вперед.

Для обеспечения прозрачности решений префетчера применяется адаптированный метод оценки относительной важности признаков, основанный на анализе градиентов и локальных линейных аппроксимациях. Процесс формирования объяснения заключается в определении конкретных исторически состоявшихся запросов и параметров контекста, которые оказали решающее влияние на принятие решения об упреждающем чтении конкретного блока.

Особое внимание при проектировании архитектуры уделено минимизации накладных расходов. Вычисление параметров объяснимости выполняется асинхронно во вспомогательных потоках ядра, что исключает блокировку основного контура обработки дисковых прерываний. Если модуль объяснимости фиксирует, что предсказание сформировано на основе изолированного шума или степень уверенности модели ниже допустимого порога, сформированная команда упреждающего чтения отменяется.

Практическая реализация и анализ результатов

Оценка эффективности разработанной методологии проводилась в изолированной тестовой среде на базе модифицированного ядра операционной системы семейства Linux. В качестве рабочей нагрузки использовались специализированные профили ввода-вывода, имитирующие работу систем управления базами данных, веб-серверов и систем обработки больших массивов неструктурированной информации.

Результаты проведенных экспериментов показали, что использование объяснимой нейросетевой модели префетчинга позволяет повысить коэффициент попадания в системный кэш при нелинейных паттернах доступа на двадцать один процент по сравнению со стандартными эвристическими алгоритмами ядра. При этом интеграция асинхронного модуля объяснимости добавила не более одного и восьми десятых процента к общей вычислительной нагрузке на центральный процессор, что является полностью допустимым показателем.

Применение методов объяснимого искусственного интеллекта позволило получить следующие результаты:

Во-первых, удалось выявить и предотвратить случаи массового ошибочного подкачивания блоков, вызываемые короткоживущими случайными сканами дискового пространства.

Во-вторых, анализ важности признаков позволил сократить глубину анализируемой истории вызовов на тридцать процентов без потери точности предсказаний, что существенно уменьшило требования модели к оперативной памяти.

В-третьих, сформированы правила автоматического переключения на классический алгоритм в условиях критической перегрузки дискового контроллера.

Заключение

Внедрение методов объяснимого искусственного интеллекта в подсистемы ввода-вывода операционных систем открывает новые возможности для повышения производительности систем хранения данных. Прозрачность и контролируемость

процессов принятия решений позволяют эффективно использовать потенциал глубокого обучения без риска деградации стабильности и отзывчивости системы.

Дальнейшие направления исследований связаны с переносом алгоритмов вычисления объяснений непосредственно в прошивку интеллектуальных твердотельных накопителей и оптимизацией работы модели в условиях распределенных сетевых файловых систем.

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.