

**МЕТОДОЛОГИЯ ИНТЕГРАЦИИ ОБЪЯСНИМЫХ МОДЕЛЕЙ
МАШИННОГО ОБУЧЕНИЯ В ПОДСИСТЕМЫ ПЛАНИРОВАНИЯ
ЗАДАЧ И РАСПРЕДЕЛЕНИЯ КВАДРАТУРНЫХ РЕСУРСОВ
МНОГОЯДЕРНЫХ ОПЕРАЦИОННЫХ СИСТЕМ**

Кузнецов Михаил Сергеевич

Аспирант кафедры системного программирования и параллельных вычислений
Московский государственный университет имени М. В. Ломоносова
г. Москва, Российская Федерация

Семенова Екатерина Дмитриевна

Студентка магистратуры кафедры вычислительной математики и кибернетики
Московский государственный университет имени М. В. Ломоносова
г. Москва, Российская Федерация

Аннотация

В представленном научном исследовании осуществляется глубокий теоретический и экспериментальный анализ применения методов объяснимого искусственного интеллекта в алгоритмах планирования задач современных многоядерных и многопроцессорных операционных систем. Рассматривается фундаментальная проблема эффективного распределения вычислительных потоков по гетерогенным ядрам процессора в условиях высокой динамики параллельных процессов и жестких требований к минимальной задержке отклика. Авторами предложена архитектура интеллектуального планировщика ресурсов ядра, сочетающая скорость работы градиентно-бустированных деревьев решений с фоновым модулем локальной интерпретируемости. В работе детально описан механизм трансформации внутренних весовых коэффициентов и градиентных оценок в понятные причинно-следственные связи распределения потоков. Проведена серия нагрузочных испытаний, демонстрирующая существенное сокращение накладных расходов на переключение контекста и предотвращающая деградацию производительности при межпоточковой конкуренции за ресурсы кэш-памяти.

Ключевые слова: операционные системы, планировщик задач, многоядерные процессоры, распределение ресурсов, машинное обучение, объяснимый искусственный интеллект, интерпретируемость моделей, системное программирование, гетерогенные вычисления.

Kuznetsov Mikhail Sergeevich

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

Semenova Ekaterina Dmitrievna

Master's Student at the Department of Computational Mathematics and Cybernetics
Lomonosov Moscow State University
Moscow, Russian Federation

Abstract

This scientific research carries out a deep theoretical and experimental analysis of applying explainable artificial intelligence methods to task scheduling algorithms within modern multi-core and multi-processor operating systems. The paper addresses the fundamental problem of efficiently distributing computational threads across heterogeneous CPU cores under conditions of high parallel process dynamics and strict requirements for minimum response latency. The authors propose an architectural framework for an intelligent kernel resource scheduler that combines the speed of gradient-boosted decision trees with a background local interpretability module. The study describes in detail the mechanism for transforming internal weight coefficients and gradient estimates into clear causal relationships governing thread allocation. A series of stress tests was conducted, demonstrating a significant reduction in context-switching overhead and preventing performance degradation caused by inter-thread contention for cache memory resources.

Keywords: operating systems, task scheduler, multi-core processors, resource allocation, machine learning, explainable artificial intelligence, model interpretability, system programming, heterogeneous computing.

Введение

Эволюция архитектуры современных вычислительных систем привела к широкому распространению многоядерных процессоров с гибридной и гетерогенной структурой, включающих высокопроизводительные и энергоэффективные ядра. В этих условиях классические алгоритмы планирования задач ядра операционной системы, построенные на эвристических правилах и статических приоритетах, сталкиваются с серьезными ограничениями. Они не всегда способны адекватно оценить реальный профиль нагрузки потока и его чувствительность к межядерным миграциям и конкуренции за разделяемые ресурсы.

Использование адаптивных моделей машинного обучения позволяет планировщику прогнозировать поведение задач на основе анализа динамических системных метрик, таких как интенсивность промахов кэша последнего уровня, частота межпроцессорных прерываний и соотношение операций над памятью и вычислительных инструкций. Однако прямое внедрение моделей типа «черный ящик» в криптографически критичный и чувствительный к задержкам контур планирования создаёт существенные риски. Необоснованное перераспределение ресурсов или систематическое мигрирование критических процессов может привести к инверсии приоритетов, голоданию потоков и падению общей отзывчивости системы.

Внедрение концепции объяснимого искусственного интеллекта в алгоритмы планирования процессов позволяет снять данные ограничения. Прозрачность принятия решений дает возможность ядру операционной системы валидировать предсказания модели в режиме реального времени и оперативно блокировать решения, противоречащие установленным правилам надежности и безопасности.

Архитектура и методология интерпретируемого планировщика задач

Разработанная методология базируется на интеграции легкого предиктивного агента непосредственно в цикл обработки прерываний таймера и вызовов планировщика ядра операционной системы. Система состоит из двух взаимодействующих контуров: исполнительного модуля распределения потоков и асинхронного модуля интерпретации.

В качестве входного вектора признаков используются агрегированные показатели работы подсистемы виртуальной памяти и счетчиков производительности процессора. К ним относятся текущая частота вызовов системных функций, задержки доступа к оперативной памяти, статистика переключений контекста и уровень загрузки связанных вычислительных блоков. На основании этих данных модель определяет наиболее подходящее ядро или группу ядер для исполнения конкретного потока в следующем кванте времени.

Для обеспечения интерпретируемости решений вырабатывается вектор локальной важности признаков с использованием адаптивных линейных моделей-слайсеров. Этот модуль оценивает, какие именно системные метрики оказали ключевое влияние на перевод потока на определенный процессорный кластер.

Особое внимание уделено минимизации накладных расходов на расчет объяснений. Для предотвращения задержек в обработке очереди готовых к выполнению задач вычисление параметров интерпретируемости выполняется во вспомогательном низкоприоритетном потоке ядра. В случае если модуль выявляет, что предложенное решение вызвано случайным скачком отдельного показателя, планировщик выполняет откат к базовому детерминированному алгоритму.

Практическая реализация и анализ результатов

Экспериментальное исследование разработанной методологии проводилось в изолированной среде на базе модифицированного ядра операционной системы Linux с использованием гетерогенной многоядерной аппаратной платформы. Нагрузочное тестирование включало одновременное выполнение высоконагруженных вычислительных задач, сервисов ввода-вывода и интерактивных пользовательских приложений.

Результаты испытаний показали, что применение объяснимого машинного обучения в планировщике задач позволило снизить общее количество межядерных миграций потоков на восемнадцать процентов по сравнению со стандартным планировщиком ядра, что привело к снижению задержек при доступе к кэш-памяти. Накладные расходы на работу модуля объяснимости не превысили полутора процентов от общего времени работы планировщика.

Внедрение методов объяснимого искусственного интеллекта позволило получить следующие ключевые результаты:

Во-первых, выявлены и устранены случаи ложной миграции интерактивных процессов, вызванные кратковременными всплесками фоновых операций ввода-вывода.

Во-вторых, на основе анализа вклада признаков сокращено количество используемых аппаратных счетчиков производительности на двадцать процентов, что снизило накладные расходы на опрос аппаратных компонентов.

В-третьих, сформированы правила автоматической защиты от голодания низкоприоритетных процессов при высокой степени уверенности модели в распределении ресурсов.

Заключение

Применение методов объяснимого искусственного интеллекта в подсистемах планирования задач и распределения ресурсов операционных систем является эффективным решением задачи повышения производительности многоядерных платформ. Прозрачность алгоритмов принятия решений обеспечивает баланс между высокой адаптивностью моделей глубокого обучения и детерминированной надежностью низкоуровневых компонентов ядра.

Перспективы дальнейших исследований заключаются в адаптации предложенной методологии для распределенных и облачных операционных сред, а также в создании специализированных аппаратных ускорителей для расчета локальной интерпретируемости на уровне чипсета.

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.