

به نام خدا



وزارت علوم، تحقیقات، و فناوری

پژوهشگاه علوم و فناوری اطلاعات ایران

(ایرانداک)

خلاصه گزارش طرح پژوهشی سازمانی

عنوان طرح پژوهشی

توسعه زیرساخت سخت افزاری هوش مصنوعی مقرون به صرفه و توسعه پذیر برای پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک)

تاریخ شروع طرح پژوهشی	۱۴۰۳/۸/۹	تاریخ پایان طرح پژوهشی	۱۴۰۴/۵/۲۱
نام و نام خانوادگی مجری طرح پژوهشی	عمار جلالی منش		
نام و نام خانوادگی همکاران طرح پژوهشی			

درباره طرح پژوهشی

این پژوهش به طراحی و پیاده سازی یک سرور هوش مصنوعی مقرون به صرفه با قابلیت پردازش موازی در پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک) می پردازد.
مقدمه و بیان مسئله:

با توجه به نیاز به زیرساخت های قدرتمند برای اجرای مدل های بزرگ هوش مصنوعی، مسئله اصلی انتخاب سخت افزارهایی بود که بتواند نیازهای محاسباتی مدل هایی مانند DeepSeek-70B و Llama Vision-90B را بدون هزینه های گزاف تأمین کنند. سرور طراحی شده بر پایه پردازش موازی با استفاده از شش کارت گرافیک NVIDIA RTX 3090، یک پردازنده ۳۲ هسته ای AMD Threadripper 3970X، و ۱۶۰ گیگابایت رم DDR4 است. این رویکرد به عنوان جایگزینی اقتصادی و منعطف برای سرورهای آماده گران قیمت انتخاب شد.

هدف های پژوهش:

هدف اصلی، توسعه یک زیرساخت سخت افزاری برای پردازش هوش مصنوعی بود که شامل طراحی و ساخت یک سرور قدرتمند با استفاده از چندین پردازنده گرافیکی به صورت موازی است. همچنین، ارائه راهکارهایی برای کاهش مشکلات مرتبط با استفاده از چندین کارت گرافیک در یک سرور، مانند مدیریت حرارت، پهنای باند، سازگاری نرم افزاری و مدیریت منابع، از دیگر اهداف پژوهش بود.

آنچه انجام شده است:

این پروژه در پنج گام اصلی انجام شد:

۱. طراحی معماری سرور و خرید تجهیزات با تمرکز بر تعادل عملکرد و هزینه.
 ۲. ساخت یک شاسی سفارشی برای نصب همزمان شش کارت گرافیک و سیستم خنک کننده.
 ۳. نصب سخت افزارها و توسعه سیستم خنک کننده مایع.
 ۴. پیکربندی نرم افزاری شامل نصب ویندوز ۱۱ و پلتفرم Ollama برای مدیریت مدل ها.
 ۵. استقرار موفقیت آمیز مدل های بزرگ هوش مصنوعی مانند DeepSeek-70B و Llama Vision-90B بر روی زیرساخت ساخته شده.
- این پروژه نشان داد که با ترکیب سخت افزارهای در دسترس و ابزارهای نرم افزاری مدرن مانند Ollama، می توان زیرساخت های پردازشی قدرتمندی را پیاده سازی کرد.

روش پژوهش

روش پژوهش در این طرح بر اساس رویکردی عملیاتی و مهندسی برای طراحی و پیاده سازی یک سرور هوش مصنوعی متمرکز بوده و شامل پنج گام اصلی است:

طراحی معماری سرور و خرید تجهیزات: در این مرحله، اجزای سخت افزاری با دقت انتخاب شدند. انتخاب ها بر تعادل میان عملکرد و هزینه تمرکز داشتند تا یک زیرساخت مقرون به صرفه ایجاد شود. اجزای اصلی شامل شش کارت گرافیک NVIDIA RTX 3090، پردازنده ۳۲ هسته ای AMD Threadripper 3970X، و ۱۶۰ گیگابایت رم DDR4 بودند.

ساخت شاسی سفارشی: یک شاسی فلزی مخصوص طراحی و ساخته شد تا بتواند همزمان شش کارت گرافیک را در خود جای دهد. این شاسی به گونه ای طراحی شد که فضای کافی برای سیستم خنک کننده هوا-مایع را نیز فراهم کند.

نصب سخت افزار و توسعه سیستم خنک کننده: قطعات خریداری شده در شاسی سفارشی قرار گرفتند. کابل های PCIe Riser برای اتصال کارت های گرافیک به مادربرد استفاده شدند و یک سیستم خنک کننده مایع برای کنترل دمای کارت ها نصب و راه اندازی شد.

پیکربندی نرم افزاری: پس از نصب سخت افزار، سیستم عامل ویندوز ۱۱ پرو به همراه درایورهای NVIDIA نصب گردید.

سپس پلتفرم Ollama برای مدیریت و اجرای مدل های هوش مصنوعی از طریق یک رابط کاربری گرافیکی (GUI) پیاده سازی شد.

استقرار و پیکربندی مدل های هوش مصنوعی: در آخرین گام، مدل های بزرگ هوش مصنوعی مانند DeepSeek-70B و Llama Vision-90B بر روی زیرساخت ایجاد شده استقرار یافتند تا عملکرد سیستم مورد آزمایش قرار گیرد.

در این پژوهش، هیچ گردآوری یا تحلیل داده ای به معنای مرسوم آن انجام نشده است. روش کار بیشتر بر اجرا، پیاده سازی و آزمایش یک سیستم فیزیکی تمرکز داشته است.

یافته‌های طرح پژوهشی

یافته‌های اصلی این طرح پژوهشی نشان می‌دهد که با ترکیب سخت‌افزارهای در دسترس و ارزان‌قیمت‌تر، می‌توان یک زیرساخت قدرتمند و مقرون‌به‌صرفه برای پردازش مدل‌های بزرگ هوش مصنوعی ایجاد کرد. این یافته‌ها به طور مستقیم به پرسش اصلی پژوهش پاسخ می‌دهند و امکان اجرای مدل‌های پیچیده را بر روی یک سرور سفارشی‌سازی شده تأیید می‌کنند.

دستاوردهای کلیدی

کارایی هزینه‌ای: این پروژه نشان داد که می‌توان با استفاده از سخت‌افزارهای موجود در بازار (به عنوان مثال، کارت‌های گرافیک NVIDIA RTX 3090 که قیمت مناسب‌تری نسبت به کارت‌های سری حرفه‌ای A دارد)، یک سرور با توان پردازشی بالا ساخت که از نظر اقتصادی، جایگزین مناسبی برای سرورهای آماده و گران‌قیمت شرکتی است.

پردازش موازی: با موفقیت شش کارت گرافیک NVIDIA RTX 3090 در یک سرور، قابلیت پردازش موازی برای مدل‌های هوش مصنوعی فراهم شد. این پیکربندی امکان اجرای مدل‌های بسیار بزرگ و پردازش سریع‌تر داده‌ها را فراهم کرد، که یک دستاورد مهم فنی محسوب می‌شود.

مدیریت مدل‌های بزرگ: با استفاده از پلتفرم Ollama و سیستم عامل ویندوز ۱۱ پرو، این سرور توانست مدل‌های بزرگی مانند DeepSeek-70B و Llama Vision-90B را با موفقیت استقرار و اجرا کند. این امر نشان می‌دهد که ابزارهای نرم‌افزاری مدرن می‌توانند به خوبی بار پردازشی را میان پردازنده‌های گرافیکی توزیع کنند و به کاربران امکان مدیریت آسان مدل‌ها را بدهند.

راهکارهای فنی: این پژوهش به صورت عملی راهکارهایی برای غلبه بر چالش‌های فنی مانند مدیریت حرارت و پهنای باند را ارائه کرد. ساخت یک شاسی سفارشی و استفاده از سیستم خنک‌کننده مایع برای کنترل دمای کارت‌های گرافیک، نشان‌دهنده یک رویکرد مؤثر در برابر مشکلات فنی است.

در نهایت، نتایج این پژوهش ثابت کرد که با رویکردی مهندسی و عملی، امکان ساخت یک سرور هوش مصنوعی قوی و منعطف با هزینه‌ای به مراتب کمتر از گزینه‌های تجاری وجود دارد، که به ویژه برای مراکز تحقیقاتی و سازمان‌هایی با بودجه محدود حائز اهمیت است.

نتیجه‌گیری و پیشنهادها

این پروژه نشان می‌دهد که ساخت یک سرور هوش مصنوعی قدرتمند با استفاده از سخت‌افزارهای در دسترس، یک راه‌حل اقتصادی و مؤثر برای پردازش مدل‌های بزرگ است. نتایج به دست آمده ثابت کرد که می‌توان با ترکیب سخت‌افزارهای با کارایی بالا و نرم‌افزارهای مدرن، زیرساختی انعطاف‌پذیر و قدرتمند ایجاد کرد.

نتیجه‌گیری

کارایی فنی و اقتصادی: این پژوهش ثابت کرد که امکان ساخت یک سرور AI با شش کارت گرافیک NVIDIA RTX

3090 برای اجرای مدل‌های پیچیده مانند DeepSeek-70B و Llama Vision-90B وجود دارد. این رویکرد، در مقایسه با سرورهای آماده شرکتی، هزینه‌ها را به طور قابل توجهی کاهش می‌دهد.

حل چالش‌های فنی: با طراحی شاسی سفارشی و پیاده‌سازی سیستم خنک‌کننده مایع، چالش‌های مربوط به مدیریت حرارت و فضا در یک سیستم چند GPU-با موفقیت حل شد.

استفاده از ابزارهای نوین: پلتفرم Ollama نقش مهمی در تسهیل مدیریت مدل‌ها و توزیع بار پردازشی میان کارت‌های گرافیک ایفا کرد و تجربه کاربری را بهبود بخشید.

پیشنهادهای

بهینه‌سازی سیستم خنک‌کننده: برای بهبود بیشتر عملکرد، می‌توان سیستم خنک‌کننده مایع را با راه‌حل‌های پیشرفته‌تر (مانند خنک‌کننده تمام برد) بهینه‌سازی کرد تا مدیریت حرارت در بارهای پردازشی بسیار بالا نیز تضمین شود.

مطالعه مقایسه‌ای: انجام یک مطالعه مقایسه‌ای با سرورهای هوش مصنوعی آماده و استاندارد موجود در بازار، می‌تواند ارزش اقتصادی و عملکردی این رویکرد را به صورت کمی تأیید کند.

توسعه نرم‌افزاری: می‌توان نرم‌افزارهای مدیریت مدل‌ها را توسعه داد تا امکانات بیشتری مانند نظارت بر عملکرد لحظه‌ای GPU و بهینه‌سازی خودکار بار پردازشی را فراهم کند.

استفاده از پردازنده‌های قوی‌تر: برای افزایش توان پردازشی و پشتیبانی از مدل‌های پیچیده‌تر، می‌توان از پردازنده‌های قوی‌تری با تعداد هسته‌های بیشتر استفاده کرد.