



Финансирано от
Европейския съюз

Следващо Поколение ЕС



План за възстановяване и
устойчивост



Република България

ПРОТОКОЛ ЗА РАЗРАБОТВАНЕ НА ЛАБОРАТОРЕН ПРОТОТИП

на софтуер

ВИЗУАЛИЗАЦИЯ НА ОПИСАНИЕТО И ПОДБОР НА ГОЛЕМИ ЕЗИКОВИ МОДЕЛИ,

планиран резултат от

Работен пакет 2. Подбор на предварително обучени големи езикови модели
по проекта „Инфраструктура за фина настройка на предварително обучени големи езикови модели“, договор No ПБУ – 55 от 12.12.2024 г. /BG-RRP-2.017-0030-C01/
от Механизма за възстановяване и устойчивост за изпълнение на инвестиция по C2I2
„Повишаване на иновационния капацитет на Българската академия на науките (БАН) в сферата на зелените и цифровите технологии“ от Плана за възстановяване и устойчивост

Протокол: IFGPT-2025-SOFTWARE-001

Версия на протокола: 1.0

Име и версия на продукта: **Визуализация на описанието version 1.0**

Описание на продукта: **Софтуер**

Ключови думи: **големи езикови модели, описание, параметри, фина настройка**

Keywords: **LLMs, description, parameters, fine-tuning**

Ръководител: **Светла Коева**

Автори: **Светла Коева, Тинко Тинчев, Ивелина Стоянова**

Текущо ниво на техническа готовност на продукта: **TRL 4**

Дата: **15.07.2025 г.**

Общо описание

Разработен е прост уебинтерфейс за представяне на описанието на събран голям брой големи езикови модели (в текущата версия са 185 модела и вариации на модели). Интерфейсът позволява търсене в данните и филтриране на моделите по дадени параметри в описанието.

Интерфейсът е подходящ за извличане на обобщена и конкретна информация за основните характеристики на големите езикови модели, проследяване на тенденции и др.

Организация на данните

Данните на описанието са представени в CSV формат и тестова извадка е достъпна за директно изтегляне: <https://ifgpt.dcl.bas.bg/llm-descriptions/>

Алтернативно, данните се съхраняват и в JSON формат и са разработени скриптове за лесно преобразуване на форматите, като се експериментира с форматите, за да се установи най-подходящият.

Необходимо оборудване и среда за разработване на прототип на визуализацията

Визуализацията се извършва посредством скрипт на javascript, който се хоства на уебстраницата на проекта. Скриптът и извадка от данните се разпространяват и през GitHub: <https://github.com/DCL-IBL/IfGPT-LLM-Description>

Процедура за разработване на прототип на визуализацията

Протоколът за разработване на прототип включва следните основни принципи:

1. Като изходни данни се използват описанията на големи езикови модели, събрани като част от проекта. За описанието на всеки модел са използвани определен брой характеристики (Приложение 1). Въведени са някои допълнителни (обобщаващи) полета (например, наред с лиценза е въведено и поле за тип лиценз, което да позволява филтриране по обобщени характеристики на лиценза – дали е свободен), позволяващи по-ефективно търсене по релевантни параметри.
2. Визуализацията показва всички големи езикови модели, освен ако не са приложени филтри към данните.
3. Филтрирането става въз основа на Формуляр за търсене, в който се посочват конкретни желани стойности (една или множество) за отделните характеристики на данните. Във формуляра са включени само основни характеристики. Ако е релевантно, в следващите версии на визуализацията може да се добави функция за разширено търсене по повече параметри. Формулярът за търсене е интуитивен и е структуриран така, че да спестява време и усилия – при празно поле (неизбрана стойност за даден параметър) по подразбиране се приема, че не се извършва филтриране по тази характеристика и др.
4. Списъкът от модели (пълнен или филтриран) съдържа само имената на моделите и броят параметри, които са основна тяхна характеристика, различаваща разновидностите на моделите. Предложена е функция за разгръщане на описанието на модела, както и обратното му прибиране (с интуитивни бутони + и –). Форматирани за линковете в описанието на моделите. В бъдеще в следващите версии може да се подобри изгледът на визуализацията.
5. Визуализацията е представена онлайн, като демонстрира основните параметри и възможности за търсене и подбор на големи езикови модели.

Приложение 1. Списък на характеристиките, използвани за описанието на големите езикови модели

Семейство, към което принадлежи моделът
Име
Абревиатура
Създател (компания, организация)
Тип лиценз (свободен, платен)
Данни за лиценза
Достъп с чат интерфейс
Лиценз за достъп с чат интерфейс
Достъп с API
Лиценз за достъп с API
Достъп за изтегляне
Лиценз за изтегляне
Дата на публикуване
Модалност (поредица от възможности: текст към текст, аудио към текст, изображение към текст и т.н.)
Размер в брой параметри
Архитектура на модела (обща информация)
Директна връзка към официалната страница на модела (ако има)
Връзка в Ollama (ако има)
Връзка в GitHub (ако има)
Връзка в HuggingFace (ако има)
Кратко описание на предназначението на модела
Кратко описание на данните за трениране
Кратко описание на начина на трениране
Многоезичен
Заложени езици
Включва ли се български (да/не)
Включва ли се български (детайли)
Информация за данните за тестване
Връзки към публикации, описващи модела
Дължина на контекста
Допълнителни методи или особености на модела
Фина настройка

История на редакциите на протокола

Версия	Дата	Автор	Промени
1.0	15.07.2025	С. Коева и екип	Първоначален протокол