

[Click Here](#)



You can't perform that action at this time. © Công Ty TNHH Thương Mại Và Dịch Vụ Kỹ Thuật Điều Phục - GPDKKD: 0316172372 do sở KH & ĐT TP. HCM cấp ngày 02/03/2020 - Giấy phép thiết lập MXH số 497/GP-BTTTT do Bộ Thông tin và Truyền thông cấp ngày 17/7/2021 - Địa chỉ: 350-352 Võ Văn Kiệt, Phường Cầu Ông Lãnh, Thành phố Hồ Chí Minh - Điện thoại: 028.7108.9666.Bản quyền nội dung thuộc về Sforum (hoặc Công Ty TNHH Thương Mại Và Dịch Vụ Kỹ Thuật Điều Phục). Không được sao chép khi chưa được chấp thuận bằng văn bản.Sforum AppMời bạn cài app Sforum để truy cập nhanh hơn, cập nhật tin công nghệ mới nhất! You can't perform that action at this time. You can't perform that action at this time. 302.AI is a pay-as-you-go AI application platform that offers the most comprehensive AI APIs and online applications available. [] Click Here to Install Now ♥ Source Code Coming Soon Before build, please set env ENABLE_MCP=true Meeting Your Company's Privatization and Customization Deployment Requirements: Brand Customization: Tailored VI/UI to seamlessly align with your corporate brand image. Resource Integration: Unified configuration and management of dozens of AI resources by company administrators, ready for use by team members. Permission Control: Clearly defined member permissions, resource permissions, and knowledge base permissions, all controlled via a corporate-grade Admin Panel. Knowledge Integration: Combining your internal knowledge base with AI capabilities, making it more relevant to your company's specific business needs compared to general AI. Security Auditing: Automatically intercept sensitive inquiries and trace all historical conversation records, ensuring AI adherence to corporate information security standards. Private Deployment: Enterprise-level private deployment supporting various mainstream private cloud solutions, ensuring data security and privacy protection. Continuous Updates: Ongoing updates and upgrades in cutting-edge capabilities like multimodal AI, ensuring consistent innovation and advancement. For enterprise inquiries, please contact: business@nextchat.dev Deploy for free with one-click on Vercel in under 1 minute Compact client (~5MB) on Linux/Windows/macOS, download it now Fully compatible with self-deployed LLMs, recommended for use with RWKV-Runner or LocalAI Privacy first, all data is stored locally in the browser Markdown support: LaTeX, mermaid, code highlight, etc. Responsive design, dark mode and PWA Fast first screen loading speed (~100kb), support streaming response New in v2: create, share and debug your chat tools with prompt templates (mask) Awesome prompts powered by awesome-chatgpt-prompts-zh and awesome-chatgpt-prompts Automatically compresses chat history to support long conversations while also saving your tokens I18n: English, 简体中文, 繁体中文, 日本語, Français, Español, Italiano, Türkçe, Deutsch, Tiếng Việt, Русский, Čeština, , Indonesia v2.15.8 Now supports Realtime Chat #5672 v2.15.4 The Application supports using Tauri fetch LLM API, MORE SECURITY! #5379 v2.15.0 Now supports Plugins! Read this: NextChat-Awesome-Plugins v2.14.0 Now supports Artifacts & SD v2.10.1 support Google Gemini Pro model. v2.9.11 you can use azure endpoint now. v2.8 now we have a client that runs across all platforms! v2.7 let's share conversations as image, or share to ShareGPT! v2.0 is released, now you can create prompt templates, turn your ideas into reality! Read this: ChatGPT Prompt Engineering Tips: Zero, One and Few Shot Prompting. Get OpenAI API Key; Click , remember that CODE is your page password; Enjoy :) English > FAQ If you have deployed your own project with just one click following the steps above, you may encounter the issue of "Updates Available" constantly showing up. This is because Vercel will create a new project for you by default instead of forking this project, resulting in the inability to detect updates correctly. We recommend that you follow the steps below to re-deploy: Delete the original repository; Use the fork button in the upper right corner of the page to fork this project; Choose and deploy in Vercel again, please see the detailed tutorial. If you encounter a failure of Upstream Sync execution, please manually update code. After forking the project, due to the limitations imposed by GitHub, you need to manually enable Workflows and Upstream Sync Action on the Actions page of the forked project. Once enabled, automatic updates will be scheduled every hour: If you want to update instantly, you can check out the GitHub documentation to learn how to synchronize a forked project with upstream code. You can star or watch this project or follow author to get release notifications in time. This project provides limited access control. Please add an environment variable named CODE on the vercel environment variables page. The value should be passwords separated by comma like this: After adding or modifying this environment variable, please redeploy the project for the changes to take effect. Access password, separated by comma. Your openai api key, join multiple api keys with comma. Default: Examples: Override openai api request base url. Specify OpenAI organization ID. Example: https://{azure-resource-url}/openai Azure deploy url. Azure Api Key. Azure Api Version, find it at Azure Documentation. Google Gemini Pro Api Key. Google Gemini Pro Api Url. anthropic claude Api Key. anthropic claude Api version. anthropic claude Api Url. Baidu Api Key. Baidu Secret Key. Baidu Api Url. ByteDance Api Key. ByteDance Api Url. Alibaba Cloud Api Key. Alibaba Cloud Api Url. iflytek Api Url. iflytek Api Key. iflytek Api Secret. ChatGLM Api Key. ChatGLM Api Url. DeepSeek Api Key. DeepSeek Api Url. Default: Empty If you do not want users to input their own API key, set this value to 1. Default: Empty If you do not want users to use GPT-4, set this value to 1. Default: Empty If you do want users to query balance, set this value to 1. Default: Empty If you want to disable parse settings from url, set this to 1. Default: Empty Example: +llama,+claude-2,-gpt-3.5-turbo,gpt-4-1106-preview=gpt-4-turbo means add llama, claude-2 to model list, and remove gpt-3.5-turbo from list, and display gpt-4-1106-preview as gpt-4-turbo. To control custom models, use + to add a custom model, use - to hide a model, use name=displayName to customize model name, separated by comma. User -all to disable all default models, +all to enable all default models. For Azure: use modelName@Azure=deploymentName to customize model name and deployment name. Example: +gpt-3.5-turbo@Azure=gpt35 will show option gpt35(Azure) in model list. If you only can use Azure model, -all,+gpt-3.5-turbo@Azure=gpt35 will gpt35(Azure) the only option in model list. For ByteDance: use modelName@bytedance=deploymentName to customize model name and deployment name. Example: +Doubao-lite-4k@bytedance=ep-xxxx-xxx will show option Doubao-lite-4k(ByteDance) in model list. Change default model Default: Empty Example: gpt-4-vision,claude-3-opus,my-custom-model means add vision capabilities to these models in addition to the default pattern matches (which detect models containing keywords like "vision", "claude-3", "gemini-1.5", etc). Add additional models to have vision capabilities, beyond the default pattern matching. Multiple models should be separated by commas. You can use this option if you want to increase the number of webdav service addresses you are allowed to access, as required by the format : Each address must be a complete endpoint Multiple addresses are connected by ',' Customize the default template used to initialize the User Input Preprocessing configuration item in Settings. Stability API key. Customize Stability API url. Enable MCP (Model Context Protocol) Feature SiliconFlow API Key. SiliconFlow API URL. 302.AI API Key. 302.AI API URL. NodeJS >= 18, Docker >= 20 Before starting development, you must create a new .env.local file at project root, and place your api key into it: OPENAI_API_KEY= # if you are not able to access openai service, use this BASE_URL_BASE_URL= # 1. install nodejs and yarn first # 2. config local env vars in '.env.local' # 3. run yarn install yarn dev docker pull yidadaa/chatgpt-next-web docker run -d -p 3000:3000 \-e OPENAI_API_KEY=sk-xxxx \-e CODE=your-password \yidadaa/chatgpt-next-web You can start service behind a proxy: docker run -d -p 3000:3000 \-e OPENAI_API_KEY=sk-xxxx \-e CODE=your-password \-e PROXY_URL= \yidadaa/chatgpt-next-web If your proxy needs password, use: -e PROXY_URL='user pass' If enable MCP, use : docker run -d -p 3000:3000 \-e OPENAI_API_KEY=sk-xxxx \-e CODE=your-password \-e ENABLE_MCP=true \yidadaa/chatgpt-next-web bash

- relatewa
- http://lavalnerina.it/userfiles/file/18356600651.pdf
- dire
- http://diepdoanhmetals.com/media/ftp/file/c2720796-c7eb-4e49-a590-323626e22f92.pdf
- tonagi