

Riduci i Costi AI del 70-92%.

92.3%

Riduzione
Costi Massima

89-93%

Su Tutti i
Modelli LLM

0%

Perdita di
Qualità

<2 min

Tempo di
Integrazione

Come Funziona

01

Deploy come Proxy

FRUX si posiziona tra la tua app e il provider AI. Zero modifiche.

02

Ottimizzazione Auto

TVP memorizza pattern, deduplica file, instrada al modello ottimale.

03

Paghi sul Risparmio

Nessun risparmio = nessun costo. Tieni sempre la maggior parte.

Performance Validata

Test reali in produzione · Dati filing USPTO · Dic 2025

Caso d'Uso

Senza FRUX

Con FRUX

Risparmio

Modello

Proiezione Risparmi Annuali

8 sessioni AI/giorno · 20 giorni/mese · Prezzi Claude Sonnet

Team	Costo Diretto	Con FRUX	Risparmio
1 Dev	\$26,285/anno	\$2,035/anno	\$24,250
10 Dev	\$262,848/anno	\$20,352/anno	\$242,496
100 Dev	\$2.63M/anno	\$203K/anno	\$2.42M

Perche FRUX

- Zero Rischi**
Nessun risparmio = nessun costo. Mai.
- Zero Modifiche**
Proxy trasparente. Plug and play.
- Model Agnostic**
Claude, GPT-4, Gemini, qualsiasi LLM.
- Brevettato**
Tecnologia TVP (USPTO #63/947,144).
- 11% Più Veloce**
Payload ridotti = rete più rapida.

Tecnologia

SLIM Protocol
Riduzione base 40-50% dei token. Open source.

Smart Routing
Selezione automatica modello per complessità.

TVP Engine
Caching semantico. Fino a 96.3% su contenuti ripetuti.

Cross-User Cache
Elabora una volta, servi tutti gli utenti.

Inizia il tuo pilota gratuito di 30 giorni

Nessun impegno · Nessun costo iniziale · Risultati garantiti

info@frux.ai

frux.pro · +39 340 706 3047