

Comparativa de modelos Opencode Go para desarrollo fullstack

Modelo	Proveedor	Código	Agéntico	Coste/ Rend.	Mejor uso	Comentario
DeepSeek V4 Pro	DeepSeek	9	9	7	Equilibrado premium	Líder en benchmarks de código; excelente agente multi-paso, pero \$15/mes y 3.450 req/5h
DeepSeek V4 Flash	DeepSeek	7	7	10	Económico para prototipado	\$0,14/M token — 31.650 req/5h; ideal para iterar rápido y tareas de alto volumen
Qwen3.7 Max	Alibaba (Qwen)	9	8	6	Código premium	Rendimiento cercano a DeepSeek V4 Pro, pero más caro (\$2,50/\$7,50); 950 req/5h
Qwen3.7 Plus	Alibaba (Qwen)	7	7	8	Equilibrado versátil	Buen balance calidad/velocidad; 4.300 req/5h, soporta 256K+ con recargo
Kimi K2.7 Code	Moonshot AI	8	6	7	Código puro	Especializado en generación y depuración; 1.350 req/5h, precio moderado (\$0,95/\$4,00)
GLM-5.2	Zhipu AI	7	7	8	Equilibrado versátil	Buena relación calidad-precio; 880 req/5h y rendimiento sólido en código bilingüe
MiniMax M3	MiniMax	7	6	7	Tareas agénticas ligeras	Capaz con herramientas básicas; 3.200 req/5h a \$0,30/\$1,20, buen throughput
Grok 4.5	xAI	8	7	5	Código premium	Razonamiento fuerte, pero caro (\$2/\$6) y solo 120 req/5h — para uso esporádico
Kimi K3	Moonshot AI	7	8	5	Tareas agénticas	Contexto ultralargo (1M+) ideal para agentes con historial extenso; solo 110 req/5h
MiMo-V2.5-Pro	—	7	6	6	Código versátil	Rendimiento sólido a precio medio (\$0,435/\$0,87); 3.250 req/5h en tier \$15
Qwen3.6 Plus	Alibaba (Qwen)	6	6	7	Económico para prototipado	Generación anterior, todavía capaz; 3.300 req/5h, adecuado para tareas estándar
MiniMax M2.7	MiniMax	6	5	8	Económico rápido	Buena opción económica con 3.400 req/5h; suficiente para autocompletado y debugging simple
GLM-5.1	Zhipu AI	6	6	8	Económico para prototipado	Similar a GLM-5.2 pero ligeramente inferior; 880 req/5h
Kimi K2.6	Moonshot AI	6	5	8	Modelo rápido de soporte	Versión anterior de K2, más barata; 1.150 req/5h para tareas auxiliares
Hy3	—	5	4	8	Modelo rápido de soporte	Modelo de bajo coste (\$0,14/\$0,58); 4.300 req/5h, útil como respaldo rápido
MiMo-V2.5	—	5	5	9	Económico para prototipado	Extremadamente barato (\$0,14/\$0,28); 30 req/5h — ideal para pruebas exploratorias

Notas: Las puntuaciones se basan en benchmarks públicos (HumanEval, MBPP, SWE-bench) y análisis de terceros disponibles a julio 2026. "Código" evalúa generación, depuración, refactorización y cobertura de lenguajes. "Agéntico" evalúa manejo de herramientas, planificación multi-paso y seguimiento de instrucciones complejas. "Costo-Rendimiento" pondera precio por token, volumen de requests incluido y calidad relativa. Proveedor "—" indica que no hay información concluyente sobre su origen. Datos de precios y límites extraídos de [opencode.ai/go](#) (julio 2026).