

Prijsstellingsstrategieën van AI-bedrijven in 2023–2026 vergeleken met Kotler

Executive summary

De meeste grote AI-aanbieders zijn in 2023–2026 geconvergeerd op **usage-based pricing** (meestal per **input- en outputtokens**) voor API-aanbod, aangevuld met **abonnementen** (B2C/B2B) en **commitment-/capaciteitsmodellen** voor voorspelbaarheid en korting (bijv. gereserveerde throughput). ¹ Dit sluit sterk aan bij de economische logica van cloud: in plaats van een eenmalige licentie is “pay-as-you-go” de standaard, met een complexe set opties en meters. ²

Vergeleken met **klassieke prijsstrategieën van Philip Kotler** (o.a. skimming/penetration, product-mix-pricing zoals bundling en captive/two-part pricing, plus prijsaanpassingen zoals kortingen en segmented pricing) zijn er duidelijke overeenkomsten: AI-bedrijven doen intensief aan **productlijn-prijsstelling** (model-tiers), **bundling** (abonnementen met meerdere features) en **(twee-delige) tariefstructuren** (vaste fee + variabel gebruik). ³

Tegelijkertijd ontstaan hybride en “cloud-native” varianten die Kotlers indeling slechts gedeeltelijk dekt: **token-metering** als micro-unit, **context-window-drempels** (prijs verandert bij >200k tokens), **caching-prijzen** (goedkopere “cached input” of “cache read”), en **service tiers** (Flex/Batch/Reserve) die vraag en latency sturen—een combinatie van prijsstrategie en capaciteitsmanagement. ⁴

Afbakening en methode

Dit rapport vergelijkt actuele prijspraktijken (primair 2023–2026) van vijf aanbieders: OpenAI ⁵, Google ⁶ (via Vertex AI/Google Cloud), Anthropic ⁷, Microsoft ⁸ (Azure OpenAI/Azure AI) en Cohere ⁹. Prijsdata zijn waar mogelijk gebaseerd op **officiële pricing- of documentatiepagina's**; waar prijzen contractueel/regio-afhankelijk zijn, wordt dat expliciet benoemd. ¹⁰

Belangrijke beperkingen: - **Valuta en regio**: prijzen kunnen lokaal worden weergegeven (bijv. CHF/EUR) of afhangen van deployment type/regionale data-zones. ¹¹

- **List price vs. enterprise**: enterprise-/volumeafspraken zijn vaak “contact sales”, waardoor kortingen niet publiek vergelijkbaar zijn. ¹²

- **Unit-definities**: “token”, “character”, “query”, “training hour” en “cached token” zijn niet volledig uitwisselbaar; dat is juist een kernpunt van de strategie. ¹³

Huidige prijsstellingspraktijken per aanbieder

Onderstaande mermaid-flow vat de dominante prijslogica samen: abonnementen en trials bestaan, maar productie-gebruik wordt vrijwel altijd gemeten en gefactureerd op verbruik (tokens/requests/compute), met differentiatie via tiers, caching en commitments.

```
flowchart LR
```

```
A[Toegang] -->|Free / Trial| B[Beperkingen: quota/rate limits]
```

```

A -->|Abonnement| C[Vaste fee per gebruiker/maand]
A -->|API pay-as-you-go| D[Metering]
C --> D
D --> E{Metriek}
E -->|Tokens in/out| F[LLM inference]
E -->|API calls / queries| G[Search/rerank/agents]
E -->|Compute-uren| H[Fine-tuning / provisioned capacity]
F --> I[Prijslevers: model-tier, contextdrempels, caching]
G --> I
H --> I
I --> J[Kortingen: batch/flex/commitments + enterprise deals]
J --> K[Factuur/chargeback]

```

14

OpenAI (API + fine-tuning + ChatGPT abonnementen)

Voor API's is het model typisch: prijs per **1M tokens**, apart voor **input**, **cached input** en **output**, met meerdere modellen en tiers. Bijvoorbeeld: *gpt-4o* \$2.50/1M input, \$1.25/1M cached input, \$10/1M output; *o3* \$2.00/1M input, \$0.50/1M cached input, \$8.00/1M output. ¹⁵ Dezelfde pagina toont ook meerdere **service tiers** (o.a. Flex) met andere tarieven, wat prijs inzet als keuzevariabele voor latency/throughput. ¹⁶

Fine-tuning introduceert extra meters: **training per 1M tokens** of zelfs **training per uur** (bij reinforcement fine-tuning). Voorbeelden: GPT-4.1 fine-tuning training \$25/1M tokens; o4-mini reinforcement fine-tuning training \$100 per training hour. ¹⁷

Aan de abonnementskant is er een expliciete freemium- en tier-ladder: ChatGPT Go (\$8/maand, US-prijs), Plus (\$20/maand) en Pro (\$200/maand). ¹⁸ Voor zakelijke teams is "Business" per gebruiker per maand (weergave is gelokaliseerd; in de bron wordt CHF 29/user/month bij jaarfacturatie getoond) en "Enterprise" is contact-sales. ¹⁹

Google (Vertex AI / Google Cloud Generative AI)

Vertex AI hanteert pay-as-you-go met **prijs per 1M tokens**, maar met een opvallende variatie: prijs verschilt bij **contextwindow-drempels** ($\leq 200k$ tokens versus $> 200k$ tokens) en er zijn expliciete tarieven voor **cached input**. Bijvoorbeeld voor Gemini 3.0 Pro (tekst): input \$1.25/1M tokens ($\leq 200k$) en \$2.50/1M tokens ($> 200k$); output \$5.00/1M tokens ($\leq 200k$) en \$10.00/1M tokens ($> 200k$); cached input is lager geprijsd.

Daarnaast bestaat "Model Optimizer" met een per-response prijs en drempel: er wordt pas gefactureerd na 200 responses; errors worden niet doorbelast.

Qua gratis instap gebruikt Google Cloud klassieke cloud-mechanismen: **\$300 trial credit** voor nieuwe klanten (met verdere free-tier opties). ²⁰

Anthropic (Claude API + Claude abonnementen)

Anthropic's API-prijsbasis is eveneens tokens in/out per 1M, met duidelijke "value ladder" tussen modellen: Claude Haiku 3.5 start bij \$0.80/1M input tokens en \$4/1M output tokens; Claude Sonnet 4 bij \$3/1M input en \$15/1M output; Claude Opus 4.5 bij \$15/1M input en \$75/1M output. ²¹ Daarbovenop documenteert Anthropic expliciet aanvullende meters zoals **prompt caching** en gerelateerde prijsregels (cache-economics).

Aan de B2C-kant is Claude Pro een klassiek abonnement: \$20/maand (US) of \$17/maand bij jaarabonnement (\$200 upfront); prijzen zijn regio-afhankelijk. ²²

Microsoft Azure AI / Azure OpenAI Service

Azure OpenAI combineert **pay-as-you-go tokens** met capaciteits-/commitment-modellen: Standard (on-demand) rekent input/output tokens; Provisioned Throughput Units (PTUs) reserveren capaciteit met maand-/jaarreserveringen om kosten te verlagen; Batch API biedt 50% korting op Global Standard pricing (met langere doorlooptijd). ²³ Azure maakt prijs bovendien expliciet stuurbaar via **deployment types** (Global / Data Zone EU-US / Regional). ²³

Concreet voorbeeld uit de Azure pricing-informatie (let op: weergave kan verschillen per regio/selectie): GPT-4o-Transcribe noemt tekst-input \$2.50/1M tokens en tekst-output \$10/1M tokens, naast audio-meters. ²⁴ Voor instap gebruikt Azure ook credits ("Try Azure for free" met \$200 credit voor 30 dagen). ²⁵

Cohere (API + trial keys + enterprise private deployment)

Cohere maakt expliciet onderscheid tussen **trial usage** en **production usage**: trial API keys zijn gratis maar beperkt; production keys zijn betaald en minder beperkt. ²⁶ Voor generatieve modellen is pricing token-gebaseerd; Command R+ (08-2024) vermeldt \$2.50/1M input tokens en \$10/1M output tokens. ²⁷ Cohere documenteert bovendien dat "billed units" (billed input/output tokens) kunnen afwijken van ruwe token-tellingen door model-/systeemgedrag, wat pricing en cost-forecasting minder triviaal maakt dan "tokens × tarief". ²⁸

Voor enterprise-behoefte (model customization, private deployment) is pricing "custom" en afhankelijk van eisen—klassiek 'contact sales'. ²⁹

Ter illustratie van alternatieve meters buiten tokens: via Amazon Web Services ³⁰ Bedrock wordt Cohere Rerank 3.5 bijvoorbeeld per **1.000 queries** geprijsd (met constraints per query/document chunk). ³¹

Vergelijkende tabel en kernverschillen

Aanbieder	Primaire prijsmodellen	Hoofdmeter(s)	Voorbeeld list price (tekst-LLM)	Gratis instap	Volume/ enterprise-mechanismen
OpenAI ⁵	API pay-as-you-go + fine-tuning tarieven + abonnementen (B2C/B2B) ³²	1M tokens (input/ cached/output), training tokens, training hours ³³	gpt-4o: \$2.50 in / \$10 out per 1M tokens (cached input goedkoper) ¹⁵	ChatGPT Free + betaalde tiers (Go/ Plus/Pro) ³⁴	Enterprise "contact sales"; business per se ¹⁹
Google ⁶	Cloud pay-as-you-go (Vertex AI) + trial credits ³⁵	1M tokens (in/out), cached tokens; soms per response (optimizer)	Gemini 3.0 Pro (≤200k): \$1.25 in / \$5 out per 1M tokens; >200k duurder	\$300 trial credit (90 dagen) + free tier (select producten) ²⁰	Enterprise cloud-contracten; prijs vaak integraal onderdeel van cloud spend management ³⁶
Anthropic ⁷	API pay-as-you-go + abonnementen (Claude Pro e.a.) ³⁷	1M input/output tokens; aanvullende meters (o.a. caching)	Sonnet 4: \$3 in / \$15 out per 1M tokens; Opus 4.5 hoger ³⁸	Claude Free; Pro \$20/mo of \$17/mo bij jaarplan ³⁹	Team/enterprise contractueel (niet altijd publiek) ⁴⁰

Aanbieder	Primaire prijsmodellen	Hoofdmeter(s)	Voorbeeld list price (tekst-LLM)	Gratis instap	Volume/ enterprise-mechanism
Microsoft ⁸	Azure OpenAI: pay-as-you-go + PTU-reserveringen + batch-korting ²³	Tokens (input/ output), throughput-capaciteit (PTU), batch jobs ²³	GPT-4o-Transcribe voorbeeld: \$2.50 tekst-input / \$10 tekst-output per 1M tokens ²⁴	Azure free trial (\$200 credit/30 dagen) ²⁵	PTU maand/ jaar-reserveringen; region/data zone keuze ²³
Cohere ⁹	API pay-as-you-go + trial keys (gratis/ beperkt) + enterprise private deployment ⁴²	1M tokens (in/out); billed units kunnen afwijken ⁴³	Command R+ (08-2024): \$2.50 in / \$10 out per 1M tokens ²⁷	Trial/ evaluation keys gratis maar beperkt ²⁶	Custom pricing voor customization/private deployment ²⁹

Kernobservatie: hoewel “\$ per 1M tokens” een ogenschijnlijk uniforme meeteenheid is, blijken de **échte strategische verschillen** te zitten in (i) welke tokens worden geteld (raw vs billed, cached vs non-cached), (ii) hoe prijs verandert met contextlengte en latency-eisen, en (iii) hoe sterk abonnementen/credits/capaciteitsreserveringen de marginale prijs dempen of juist escaleren. ⁴⁵

Kotler-mapping en waar het schuurt

Kotlers klassieke prijsstrategieën als referentiekader

Kotler's (en Kotler-gerelateerde Pearson-literatuur) onderscheidt onder meer: - **New-product pricing:** *market-skimming vs market-penetration*. ⁴⁶

- **Product-mix pricing:** product line, optional product, **captive product** en **product bundle**; bij diensten vaak **two-part pricing** (vaste fee + variabele usage rate). ⁴⁷

- **Price-adjustment strategies:** *discount and allowance*, **segmented pricing** (prijsdiscriminatie) en **psychological pricing**. ⁴⁶

Deze categorieën zijn historisch product-/marktgericht. Cloud- en SaaS-literatuur laat zien dat digitale diensten in de praktijk vaak juist worden geprijsd als **variabele kostenmodellen** met veel opties (pay-as-you-go, reserveringen, etc.). ⁴⁸

Expliciete overeenkomsten: “AI pricing” als Kotler-productmix + cloud-metering

Product line pricing (Kotler) ↔ model-tiering / capability tiering

Bij alle aanbieders bestaan meerdere modellen (mini/flash/haiku vs pro/opus/flagship) met sterke prijsstappen. Dit is textbook *product line pricing* (verschillende prijsstappen binnen één productlijn), maar hier gemeten in tokens in/out. ⁴⁹

Product bundle pricing (Kotler) ↔ abonnementen die meerdere features bundelen

Chat/assistants worden gebundeld met extra tools (analysis, voice, agents, admin controls) tegen een per-user prijs, waardoor de marginale tokenprijs impliciet wordt. Zowel OpenAI Business/Enterprise als Claude Pro zijn hier voorbeelden. ⁵⁰

Captive product pricing / two-part pricing (Kotler) ↔ vaste fee + variabele meters / credits

Het klassieke “razor-and-blades” idee (laag geprijsd basisproduct, hoger geprijsde verbruikscomponent) verschijnt in AI als: (a) abonnement per seat plus (b) extra credits/usage voor zwaardere modellen of enterprise-features; bij diensten is dit precies Kotlers “two-part pricing”. ⁵¹

Segmented pricing (Kotler) ↔ price discrimination via tiers, regio's, latency en commitments

Prijsdifferentiatie gebeurt niet alleen via “klantsegmenten” maar via: deployment type (regional vs global/data zone), reserved throughput (PTU), batch-kortingen en flex-tiers. Dit is functioneel prijsdiscriminatie op basis van service level en contractduur. ⁵²

Skimming vs penetration (Kotler) ↔ premium frontier models + goedkope instapmodellen/free tiers

De ladder “Haiku/Flash/mini” versus “Opus/Pro/Thinking” en daarnaast trials/free tiers, past op een skimming/penetration-hybride: hoge marges op topsegmenten, brede adoptie via goedkope of gratis instap. ⁵³

Waar Kotler minder volstaat: nieuwe of hybride modellen

De onderstaande mermaid-mapping toont waar klassieke labels te grof worden: veel moderne AI-pricing combineert meerdere Kotler-strategieën én voegt “cloud-mechanismen” toe die tegelijk prijsstrategie en capaciteitssturing zijn.

flowchart TB

```
K1[Kotler: product line pricing] --> A1[Model-tiers: haiku/mini vs opus/pro]
K2[Kotler: bundle pricing] --> A2[Abonnementen: tools + modellen + governance]
K3[Kotler: captive/two-part pricing] --> A3[Seat fee + variabele tokens/credits]
K4[Kotler: segmented pricing] --> A4[Tiers: batch/flex/priority + regio/data zone + commitments]
K5[Kotler: skimming/penetration] --> A5[Premium frontier + free/trial instap]
A6["Nieuwe variant: token- & context-metering"] --> A4
A7["Nieuwe variant: caching-tarieven"] --> A3
A8["Nieuwe variant: optimizer (charged after N)"] --> A4
```

⁵⁴

Wat Kotler conceptueel minder scherp maakt in deze markt: - **De meeteenheid is niet “product” maar “compute-proxy”**: tokens/cached tokens/context buckets zijn operationele proxies voor inference-kosten en capacity. Cloud-economics benadrukt dat cloud-diensten pay-as-you-go en multi-option zijn; AI-pricing erft dat patroon. ⁵⁵

- **Prijs als vraag- en latency-sturing**: Flex/Batch/Reserved/Provisioned tiers zijn geen simpele korting, maar sturen doorlooptijd en resource-allocatie (een mengvorm van pricing en operations). ⁵⁶

- **“Freemium” is vaak credit-gedreven i.p.v. permanent gratis**: cloud-trials (\$300) en trial keys zijn economisch anders dan een blijvende freemium-laag; ze zijn acquisitie- en leerinstrumenten met harde grenzen. ⁵⁷

- **Marginale prijs is context-afhankelijk**: bij Vertex AI verandert de prijs bij >200k tokens; bij meerdere aanbieders bestaan cache-tarieven die de marginale prijs bewust verlagen voor hergebruik—een expliciete prikkel tot architectuurkeuzes (prompt reuse, caching). ⁵⁸

Casestudies

Casus OpenAI: “service tiers” als dynamische prijsdiscriminatie op latency

OpenAI's pricing-pagina laat zien dat dezelfde (of vergelijkbare) modellen in meerdere tiers kunnen bestaan (bijv. “Flex” met lagere tarieven). Dit lijkt op Kotler's *segmented pricing*, maar het segment is hier: **service level** (wachtijd/priority) in plaats van demografie. ⁵⁹ In dezelfde prijstabellen is “cached input” substantieel goedkoper, waardoor productarchitectuur (cache-geschikte prompts) een economisch ontwerpdoel wordt—iets wat Kotler wel in algemene zin als “price adjustment” kan vangen, maar niet op micro-meter niveau. ⁶⁰

Casus Google Vertex AI: context-drempels en caching als price-for-constraints

Vertex AI publiceert expliciet verschillende tarieven voor $\leq 200k$ versus $>200k$ tokens, plus aparte cached-input prijzen. Dit is geen klassieke “bundel” of “discount” maar een prijsmechanisme dat de echte kostendrijver (extreme contextlengte) internaliseert. In Kotler-termen is dit te plaatsen onder product line/price adjustment, maar het is in feite een **non-lineaire tariefstructuur** (marginale prijs stijgt na een drempel) die typisch is voor cloud-pricing. ⁶¹

Casus Cohere: trial-vs-production keys en billed_units

Cohere maakt pricing-governance expliciet: trial API key usage is gratis maar beperkt, production keys zijn betaald; daarnaast kan billing plaatsvinden op “billed input/output tokens” die afwijken van ruwe tokens doordat Cohere soms tokens “onder de motorkap” toevoegt of anders telt. ²⁶ Dit is praktisch een moderne variant van freemium + two-part pricing: gratis evaluatie, daarna verbruiksfacturatie; maar met een additionele complexiteit (billing model) die forecasting en chargeback beïnvloedt. ⁶²

Casus Azure OpenAI: PTU-reserveringen en batch-korting als cloud-commitment strategie

Azure OpenAI positioneert drie paden: Standard pay-as-you-go tokens, Provisioned Throughput Units (PTUs) met maand/jaar-reserveringen om spend te verlagen, en Batch API met 50% korting (met langere doorlooptijd). ²³ Dit is klassiek cloud-mechanisme (commitment/slotting), en past in Kotler deels onder *discount/allowance* en *segmented pricing*, maar het is tegelijk een operationeel contractmechanisme (capaciteit reserveren). ⁶³

Conclusies: hybride modellen, implicaties en “Kotler-plus”

De actuele prijsstrategieën van AI-bedrijven zijn het best te begrijpen als een **hybride van Kotler's product-mix-denken en cloud/SaaS-metering**. Kotler blijft nuttig voor het “waarom” (segmenten, bundels, productlijnen), maar schiet tekort in het “hoe” van moderne digitale metering: tokens, caching, context-drempels en latency-tiers zijn fijnmaziger dan klassieke retail-/FMCG-analogieën. ⁶⁴

Een bruikbaar “Kotler-plus” raamwerk voor 2023–2026 is daarom: - **Product line pricing → model-/capability-versioning** (meerdere modellen met oplopende waarde en oplopende \$/token). ⁶⁵

- **Two-part / multi-part tariffs → seat fee + inbegrepen quota + overage** (abbonementen en enterprise credits) én parallel pay-as-you-go API. ⁵¹

- **Segmented pricing → service-level, regio en commitment-gebaseerde discriminatie** (Flex/Batch/Reserved/PTU, data zones). ⁶⁶

- **Freemium → credit/trial-gedreven acquisition** (trial keys, \$300 cloud credits) in plaats van “gratis voor altijd”, omdat inferencekosten structureel variabel zijn. ⁶⁷

Praktische implicatie voor afnemers (FinOps/inkoop/architectuur): de grootste kostenverschillen komen zelden uit “welke aanbieder is het goedkoopst per token” alleen, maar uit **architectuur-compatibiliteit**

met prijslevers (caching, batch, context-discipline), plus contractkeuzes (commitments/PTU) en governance (chargeback op billed units). ⁶⁸

¹ ⁹ ¹⁰ ¹⁵ ¹⁶ ³² ³³ ⁴⁵ ⁴⁹ ⁵⁶ ⁵⁸ ⁵⁹ ⁶⁰ <https://developers.openai.com/api/docs/pricing/>
<https://developers.openai.com/api/docs/pricing/>

² ¹⁴ ⁴⁸ ⁵⁵ ⁶¹ "The Economics of the Cloud"
https://www.tse-fr.eu/sites/default/files/TSE/documents/doc/wp/2024/wp_tse_1520.pdf?utm_source=chatgpt.com

³ ⁴⁶ ⁴⁷ ⁵⁴ ⁶³ ⁶⁴ <https://www.pearson.com/content/dam/one-dot-com/one-dot-com/hei-samples/marketing/9781488626203-toc.pdf>
<https://www.pearson.com/content/dam/one-dot-com/one-dot-com/hei-samples/marketing/9781488626203-toc.pdf>

⁴ ¹³ <https://cloud.google.com/vertex-ai/generative-ai/pricing>
<https://cloud.google.com/vertex-ai/generative-ai/pricing>

⁵ ²⁷ <https://docs.cohere.com/docs/command-r-plus>
<https://docs.cohere.com/docs/command-r-plus>

⁶ ³⁷ ³⁸ ⁴¹ <https://docs.cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-5-pro>
<https://docs.cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-5-pro>

⁷ ⁸ ²³ ²⁴ ²⁵ ³⁰ ⁵² ⁶⁶ ⁶⁸ <https://azure.microsoft.com/en-us/pricing/details/azure-openai/>
<https://azure.microsoft.com/en-us/pricing/details/azure-openai/>

¹¹ ¹² ¹⁹ ⁵⁰ <https://openai.com/business/chatgpt-pricing/>
<https://openai.com/business/chatgpt-pricing/>

¹⁷ <https://openai.com/api/pricing/>
<https://openai.com/api/pricing/>

¹⁸ <https://openai.com/index/introducing-chatgpt-go/>
<https://openai.com/index/introducing-chatgpt-go/>

²⁰ ³⁵ ⁵⁷ ⁶⁷ Free Trial and Free Tier Services and Products
https://cloud.google.com/free?utm_source=chatgpt.com

²¹ <https://developers.openai.com/api/docs/models/gpt-4.1>
<https://developers.openai.com/api/docs/models/gpt-4.1>

²² ³⁹ Plans & Pricing | Claude by Anthropic
https://claude.com/pricing?utm_source=chatgpt.com

²⁶ ²⁸ ⁴² ⁴³ ⁴⁴ ⁶² <https://docs.cohere.com/docs/how-does-cohere-pricing-work>
<https://docs.cohere.com/docs/how-does-cohere-pricing-work>

²⁹ <https://cohere.com/pricing>
<https://cohere.com/pricing>

³¹ <https://aws.amazon.com/bedrock/pricing/>
<https://aws.amazon.com/bedrock/pricing/>

³⁴ <https://chatgpt.com/pricing/>
<https://chatgpt.com/pricing/>

³⁶ <https://azure.microsoft.com/en-us/pricing>
<https://azure.microsoft.com/en-us/pricing>

40 Anthropic intensifies AI competition with \$200 plan for Claude models

https://www.reuters.com/technology/artificial-intelligence/anthropic-intensifies-ai-competition-with-200-plan-claude-models-2025-04-09/?utm_source=chatgpt.com

51 [https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/](https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/Principles%20of%20Marketing%2014th%20Edition%20-%20Kotler%20and%20Amstrong%20%282012%29/Chapter%2011%20-%20Pricing%20Strategies.pdf)

[Principles%20of%20Marketing%2014th%20Edition%20-](https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/Principles%20of%20Marketing%2014th%20Edition%20-%20Kotler%20and%20Amstrong%20%282012%29/Chapter%2011%20-%20Pricing%20Strategies.pdf)

[%20Kotler%20and%20Amstrong%20%282012%29/Chapter%2011%20-%20Pricing%20Strategies.pdf](https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/Principles%20of%20Marketing%2014th%20Edition%20-%20Kotler%20and%20Amstrong%20%282012%29/Chapter%2011%20-%20Pricing%20Strategies.pdf)

[https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/](https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/Principles%20of%20Marketing%2014th%20Edition%20-%20Kotler%20and%20Amstrong%20%282012%29/Chapter%2011%20-%20Pricing%20Strategies.pdf)

[Principles%20of%20Marketing%2014th%20Edition%20-%20Kotler%20and%20Amstrong%20%282012%29/](https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/Principles%20of%20Marketing%2014th%20Edition%20-%20Kotler%20and%20Amstrong%20%282012%29/Chapter%2011%20-%20Pricing%20Strategies.pdf)

[Chapter%2011%20-%20Pricing%20Strategies.pdf](https://nscpolteksby.ac.id/ebook/files/Ebook/Business%20Administration/Principles%20of%20Marketing%2014th%20Edition%20-%20Kotler%20and%20Amstrong%20%282012%29/Chapter%2011%20-%20Pricing%20Strategies.pdf)

53 65 <https://www.anthropic.com/claude/haiku>

<https://www.anthropic.com/claude/haiku>