

Limites e Potenciais da Inteligência Artificial na Avaliação de Redações do Enem: Um Estudo Comparativo entre Correção Automatizada e Mentoria Especializada

Cristiane Nordi¹

¹Universidade Federal de São Carlos, São Carlos, Brasil, (cris.nordi@hotmail.com)

Resumo: Este estudo compara a correção de uma redação do Enem por duas IAs generativas (ChatGPT e Gemini) e por uma mentora humana. As IAs destacam-se na análise microestrutural (gramática, coesão), mas falham na avaliação macroestrutural (repertório, projeto de texto) e chegam a confabular repertórios inexistentes no texto original. Conclui-se que a IA atua de forma complementar, sendo a mentoria humana indispensável para a autoria crítica e a emancipação intelectual do estudante.

Palavras-chave: Inteligência Artificial; Redação do Enem; Mentoria de Escrita; Avaliação Formativa; Alucinação em Modelos de Linguagem; Tecnologias na Educação.

INTRODUÇÃO

A Complexidade da Grade de Avaliação do Enem e as Competências do Inep. A redação do Enem exige do candidato a produção de um texto dissertativo-argumentativo sobre um tema de ordem social, científica, cultural ou política, no qual ele deve defender um ponto de vista claro e propor uma intervenção social. A avaliação oficial, coordenada pelo Instituto Nacional de Estudos e Pesquisas Educacionais Anísio Teixeira (INEP), estrutura-se em cinco competências básicas que pontuam de 0 a 200 pontos cada, totalizando a nota máxima de 1000 pontos. A Competência 1 avalia o domínio da norma culta da língua escrita; a Competência 2 foca na compreensão da proposta, na aplicação das várias áreas do conhecimento (repertório sociocultural) e nos limites estruturais

do texto; a Competência 3 analisa a seleção, relação, organização e interpretação de informações (projeto de texto e argumentação); a Competência 4 avalia o conhecimento dos mecanismos linguísticos de coesão; e a Competência 5 julga a elaboração de uma proposta de intervenção articulada que respeite os direitos humanos. Conforme os manuais de correção do Inep, enquanto as Competências 1 e 4 se concentram prioritariamente em aspectos formais e microestruturais da língua, as Competências 2, 3 e 5 exigem do avaliador uma leitura holística, interpretativa e profundamente hermenêutica; a Competência 3, por exemplo, não avalia a quantidade de argumentos, mas sim a presença de um "projeto de texto" estratégico e a ausência de lacunas argumentativas ou saltos lógicos (INEP, 2023). Esses conceitos são particularmente desafiadores para modelos de linguagem de grande porte, que são fundamentalmente sistemas de reconhecimento estatístico de padrões textuais, sem acesso a uma representação de mundo que sustente a compreensão semântica profunda das relações de causa e efeito em uma sociedade complexa (BENDER et al., 2021).

Avaliação Formativa versus Avaliação Somativa no Contexto Digital

A teoria da avaliação educacional distingue tradicionalmente as práticas somativas das práticas formativas. De acordo com os estudos de Hoffmann (2001) e Hadji (2001), a avaliação formativa pressupõe uma ação mediadora e processual, em que o erro não é punido com uma nota estática, mas compreendido como um degrau necessário no processo contínuo de aprendizagem, exigindo um feedback reconstrutivo, detalhado e individualizado por parte do educador. No ensino de redação mediado por tecnologias, as plataformas de correção automatizada baseadas estritamente em inteligência artificial generativa tendem a realizar uma avaliação predominantemente somativa ou puramente prescritiva. Os modelos matemáticos apontam o desvio sintático superficial, sugerem uma substituição vocabular imediata e atribuem uma nota estática e fria. Em contrapartida, a mentoria individualizada humana atua na perspectiva verdadeiramente formativa.

O mentor digital não apenas mensura metricamente o desempenho atual do estudante, mas estabelece um diálogo pedagógico aberto que acolhe as fragilidades do aluno, orienta o processo de reescrita e oferece caminhos práticos e personalizados para a reestruturação argumentativa, promovendo a autonomia intelectual do indivíduo.

Correção Automatizada por Inteligência Artificial: Potenciais e Riscos de Confabulação

A automatização da correção de textos dissertativo-argumentativos não é um fenômeno recente nem exclusivo dos LLMs generativos. Desde a década de 1960, sistemas de Automated Essay Scoring (AES) buscam operacionalizar features linguísticas mensuráveis — extensão de período, riqueza vocabular, marcadores de coesão — como proxies estatísticos da qualidade textual (SHERMIS; BURSTEIN, 2013). O que distingue a geração atual de ferramentas como ChatGPT e Gemini é a capacidade de produzir, além de uma nota, uma justificativa textual fluente que simula o raciocínio de um avaliador humano — o que confere a essas respostas uma aparência de autoridade interpretativa que nem sempre corresponde à sua confiabilidade real. Estudos recentes sobre o uso de LLMs para correção de redações em contexto de aprendizagem de língua têm demonstrado inconsistências relevantes entre modelos e uma tendência à instabilidade da nota atribuída a um mesmo texto (PACK; BARRETT; ESCALANTE, 2024), o que reforça a necessidade de cautela metodológica antes de se adotar essas ferramentas como substitutas da correção humana em contextos de alto impacto, como um exame de larga escala. Um segundo risco, distinto da mera leniência ou complacência na pontuação, é o da confabulação (também chamada de alucinação): a geração de conteúdo fluente, coerente e gramaticalmente correto, porém factualmente incorreto ou não sustentado pelo texto-fonte (RAWTE et al., 2023). Enquanto a literatura sobre alucinação em LLMs costuma investigar domínios como medicina, direito e produção de referências bibliográficas, o presente estudo revela, na seção de Resultados e Discussão, que esse mesmo fenômeno pode se manifestar na correção automatizada de redações, quando a IA atribui ao

estudante elementos textuais — como autores e repertórios — que não constam do texto original.

O Conceito de "Repertório de Bolso" na Matriz do Enem

Um conceito central para a análise da Competência 2 é o de "repertório de bolso", também chamado pelos próprios candidatos de repertório coringa. Trata-se de referências genéricas e memorizadas — citações de autores célebres, conceitos filosóficos ou sociológicos amplamente conhecidos — mobilizadas de forma automática e descontextualizada, sem relação de fato produtiva com o tema específico proposto na prova. O próprio Inep passou a alertar explicitamente sobre esse tipo de recurso na Cartilha de Redação do Enem, esclarecendo que a simples menção a um autor ou obra célebre, sem contextualização e sem articulação clara com os argumentos do texto, é classificada como repertório improdutivo e pode reduzir a nota da Competência 2, ainda que a citação seja, em si, verdadeira e pertinente em termos gerais (INEP, 2025). Esse conceito orienta a leitura crítica que a mentora humana realiza no estudo de caso a seguir, e que as duas ferramentas de IA generativa, como se verá, não foram capazes de reproduzir.

MATERIAL E MÉTODOS

Este trabalho caracteriza-se metodologicamente como um estudo de caso descritivo e analítico, adotando uma abordagem mista (qualitativa e comparativa). O procedimento de pesquisa estruturou-se de forma sistemática em três grandes pilares metodológicos de avaliação cruzada (triangulação de dados) aplicada sobre um único objeto textual:

- **O Corpus Textual de Análise:** Selecionou-se uma redação padrão dissertativo-argumentativa no modelo Enem produzida de forma manuscrita por um estudante real, cujo tema abordado foi "Caminhos para enfrentar a normalização dos jogos de azar feita pela publicidade no Brasil". O texto original foi preservado integralmente em sua estrutura sintática e vocabular, incluindo todas as suas rasuras, desvios ortográficos e escolhas argumentativas. A redação inicia com a citação da filósofa francesa Simone de

Beauvoir ("O mais escandaloso nos escândalos, é que nos habituamos a eles") e articula em seu desenvolvimento discussões com base nas falas de Caco Barcellos e de uma figura histórica citada como "Junqueira Gerra".

• **O Procedimento de Avaliação Automatizada (IAs Generativas):** O texto original foi transcrito e submetido à correção independente de duas das principais ferramentas de Inteligência Artificial generativa baseadas em Modelos de Linguagem de Grande Porte (LLMs) de ampla utilização comercial e acadêmica no mundo. Para a IA 1 (ChatGPT, versão GPT-5.5, da OpenAI) e para a IA 2 (Gemini, versão 3.5 Flash), utilizou-se um comando (prompt) idêntico e direto — reproduzido na íntegra no Apêndice D —, solicitando que a tecnologia atuasse como corretor oficial do Enem, atribuindo nota de 0 a 200 para cada uma das cinco competências e justificando a pontuação.

• **O Procedimento de Avaliação por Mentoria Humana Especializada:** Paralelamente, o mesmo texto foi submetido à avaliação cega de uma mentora profissional de escrita atuante em uma plataforma digital de aprendizagem de grande alcance nacional — ou seja, a mentora aplicou seus critérios de correção sem conhecimento prévio das notas ou das justificativas produzidas pelas duas ferramentas de IA. A mentora aplicou os mesmos critérios oficiais de correção do Inep, contudo, realizou anotações cirúrgicas e comentários pedagógicos detalhados diretamente nas linhas do texto do aluno, focando no desenvolvimento cognitivo e na reestruturação textual. A combinação desses três procedimentos permitiu confrontar não apenas a nota numérica atribuída por cada avaliador em cada competência, mas também a natureza qualitativa do erro identificado — distinguindo desvios formais, saltos argumentativos, imprecisões factuais e, como se verá adiante, casos de confabulação de conteúdo inexistente no texto-fonte. O prompt integral utilizado para instruir as duas ferramentas de IA generativa encontra-se reproduzido no Apêndice D, permitindo a réplica do procedimento. Reconhece-se, como limitação metodológica, que a data exata de acesso às ferramentas de IA não foi registrada no momento da coleta de dados, o que restringe parcialmente a

replicabilidade exata do procedimento; recomenda-se que estudos futuros documentem sistematicamente essa informação.

RESULTADOS E DISCUSSÃO

O confronto analítico entre as três avaliações revelou divergências metodológicas, quantitativas e conceituais profundas, que põem em xeque a autonomia completa dos algoritmos na avaliação de competências humanas complexas. A Tabela 1 apresenta de forma comparativa a distribuição das notas atribuídas por cada um dos avaliadores em cada competência da grade oficial. Nas subseções seguintes, discute-se cada bloco de competências à luz das justificativas textuais completas produzidas pela mentora e pelas duas IAs, reproduzidas integralmente nos Apêndices A, B e C.

Tabela 1. Comparativo das avaliações

Competência Avaliada	Nota Mentora Humana	Nota IA 1 (ChatGPT, GPT-5.5)	Nota IA 2 (Gemini, Flash) 3.5
C1 (Norma Culta)	120	160	160
C2 (Tema e Repertório)	120	200	200
C3 (Projeto e Argumentação)	160	180	200
C4 (Coesão Textual)	120	180	200
C5 (Proposta de Intervenção)	200	180	200
Nota Final	720 / 1000	900 / 1000	960 / 1000

Acumulada			
------------------	--	--	--

Análise Microestrutural e Gramatical: Competências 1 e 4

Os dados quantitativos expressos na Tabela 1 evidenciam uma acentuada discrepância na Competência 1 e na Competência 4. Enquanto a mentora humana penalizou rigidamente o texto do aluno em ambas as competências, atribuindo nota 120, as duas ferramentas de IA demonstraram uma complacência estatística alarmante, conferindo notas significativamente superiores (160 pela IA 1 e até notas máximas de 200 pela IA 2). No campo da Competência 1, a mentora identificou uma série de falhas sintáticas graves que comprometem a fluidez da leitura, destacando o desvio crasso de pontuação logo na abertura do texto, onde o estudante separou o sujeito do verbo por meio de uma vírgula imprópria: "O mais escandaloso nos escândalos, é que nos habituamos a eles". Adicionalmente, a mentora registrou falhas recorrentes de regência verbal e nominal, desvios no uso da crase ("vícios a população") e erros graves de concordância verbal ("a normalização desses jogos acarretem"). No entanto, as ferramentas de IA falharam gravemente em penalizar tais desvios de forma proporcional. Ambas justificaram suas notas alegando textualmente que o aluno possui um "vocabulário rico", "linguagem predominantemente formal" e "excelente domínio de estruturas sintáticas complexas". Essa falha algorítmica ocorre porque as IAs generativas operam calculando probabilidades de ocorrência de tokens e sequências lexicais associadas a padrões formais de escrita, e não a partir de uma compreensão semântica das normas gramaticais em jogo (BENDER et al., 2021). Elas identificam com facilidade termos sofisticados isolados, mas demonstram incapacidade técnica para mapear o truncamento sintático de períodos excessivamente longos e a quebra de regras gramaticais sutis de pontuação e regência em textos manuscritos transcritos. Na Competência 4, a IA 2 atribuiu nota máxima (200), exaltando a presença superficial de conectivos interparágrafos comuns ("Em primeira análise", "Adicionalmente", "Portanto"). O olhar cirúrgico da mentora humana, por sua vez, detectou que a repetição

viciosa de palavras estruturais como "normalização", "mídia" e "jogos de azar" gerou uma circularidade e repetitividade textual crônica, o que degrada substancialmente a qualidade da coesão intraparágrafo e justifica a nota 120.

Análise Macroestrutural, Repertório e Autoria: Competências 2 e 3

As maiores e mais preocupantes divergências metodológicas e conceituais residem no eixo avaliativo da inteligência argumentativa e do repertório acadêmico (Competências 2 e 3). A IA 1 e a IA 2 conferiram nota máxima de 200 pontos na Competência 2, elogiando amplamente a mobilização de repertórios socioculturais ao longo do texto. Contudo, a avaliação realizada pela mentora humana desconstruiu completamente essa suposta excelência textual, reduzindo a nota da Competência 2 para 120 pontos. A mentora expôs que o estudante recorreu sistematicamente ao que a Cartilha do Inep classifica como "repertório de bolso" (INEP, 2025) — isto é, citações coringas e frases decoradas de caráter ultra-genérico que podem ser facilmente adaptadas de forma artificial a praticamente qualquer problemática social, como a afirmação de Simone de Beauvoir sobre o hábito diante dos escândalos. Esse tipo de recurso adorna superficialmente o texto, mas não demonstra produtividade real nem articulação crítica intrínseca com o tema específico dos jogos de azar e da publicidade abusiva no Brasil. O ponto mais crítico da limitação analítica da inteligência artificial manifestou-se na validação do repertório histórico do segundo parágrafo de desenvolvimento. O aluno citou a frase "a escola é a única alavanca capaz de elevar ao povo ao nível da moral" atribuindo-a ao político "Junqueira Guerra". A mentora humana prontamente detectou e apontou dois erros severos de conteúdo acadêmico que minam a credibilidade da argumentação: primeiro, a grafia incorreta e distorcida do nome da figura histórica (o correto é Guerra Junqueiro); segundo, o fato de que ele é amplamente reconhecido e consagrado na história da literatura como um proeminente poeta português, e não prioritariamente como uma figura política aplicável a esse contexto sociopolítico brasileiro. Além disso, a mentora apontou que a ausência de orientação específica da escola no texto gerava um salto argumentativo bizarro, conectando de forma abrupta a falta de suporte

escolar à criação de "futuros publicitários através das redes sociais". Surpreendentemente, ambas as Inteligências Artificiais falharam em detectar o erro factual histórico, a grafia errônea do nome do autor e a total falta de produtividade lógica e interpretativa do argumento. Como as IAs operam sob princípios probabilísticos e estatísticos de distribuição lexical, e não sob uma avaliação real do conteúdo proposicional do texto (BENDER et al., 2021), a mera justaposição de palavras com forte apelo acadêmico (como "filósofa", "moral", "escola", "crítica") induz o algoritmo a deduzir erroneamente que o texto possui alta consistência intelectual e um projeto de texto estratégico bem delimitado (C3), camuflando graves lacunas reflexivas e ideias completamente inacabadas que comprometem o rigor do texto.

Confabulação e Fabricação de Repertórios: um Achado Adicional

A leitura integral das devolutivas completas (Apêndices B e C) revela um fenômeno qualitativamente distinto — e mais grave — do que a simples leniência avaliativa discutida acima: a confabulação de conteúdo inexistente no texto-fonte. Na sua justificativa para a Competência 2 (Apêndice B), a IA 1 (ChatGPT) lista como repertórios socioculturais mobilizados pelo estudante "Simone de Beauvoir; Lúcia Santaella; Paulo Freire; Programa Nacional Contra Publicidade de Jogos de Azar". Ocorre que Lúcia Santaella e Paulo Freire não aparecem em nenhum momento da redação transcrita (Apêndice A): são autores inexistentes no corpus analisado, atribuídos ao aluno pela própria IA. Não se trata, portanto, de um erro de leniência de nota — é a fabricação de dado factual sobre o próprio objeto avaliado, um comportamento consistente com o que a literatura de processamento de linguagem natural define como alucinação, isto é, a produção de conteúdo fluente e coerente, porém não sustentado pela evidência textual disponível (RAWTE et al., 2023). Um segundo indício do mesmo fenômeno aparece na devolutiva da IA 2 (Apêndice C): ao transcrever e comentar o repertório histórico do segundo parágrafo de desenvolvimento, o Gemini refere-se ao autor citado pelo aluno como "Junqueira Freire" — uma terceira variação do nome, diferente tanto do erro original do estudante ("Junqueira Gerra") quanto do nome correto (Guerra

Junqueiro), e que parece contaminada pela proximidade lexical com "Paulo Freire". Ou seja, a IA não apenas deixou de corrigir o erro do aluno: produziu uma nova informação incorreta, não presente no texto avaliado. Esses dois achados indicam que a fragilidade das ferramentas de IA generativa na correção de redações não se limita a uma tendência de notas infladas ou a uma leitura pouco crítica do repertório — ela também compromete a fidelidade factual básica do relatório de correção, o que tem implicações sérias para qualquer uso dessas ferramentas em contextos de avaliação de alto impacto, nos quais a auditabilidade e a rastreabilidade da justificativa avaliativa são indispensáveis.

Análise da Proposta de Intervenção: Competência 5

No julgamento da proposta de intervenção social (Competência 5), observou-se um fenômeno inverso de rigidez analítica. A mentora humana avaliou a proposta do estudante como plenamente articulada e metodologicamente impecável perante as exigências do Inep, atribuindo a nota máxima de 200 pontos. A mentora mapeou de forma clara e legítima a presença inequívoca de todos os cinco elementos obrigatórios exigidos pela banca examinadora: o agente (Governo Federal), a ação (criar o Programa Nacional Contra a Divulgação de Jogos de Azar), o modo/meio de execução (por meio de um decreto), o efeito pretendido (com a finalidade de vetar novas divulgações perante multa) e o detalhamento robusto (explicitando as ações subsequentes do programa por meio de leis a serem votadas e regulamentadas no congresso). Todavia, a IA 1 (ChatGPT) penalizou injustamente a proposta do aluno, reduzindo a nota para 180 pontos sob a justificativa textual de que a proposta "ficou um pouco genérica" e que "faltou detalhar melhor o modo de execução e o funcionamento prático das campanhas de fiscalização". Esse erro de julgamento da IA ocorre devido à sua incapacidade de realizar uma leitura estrutural e funcional integrada do parágrafo conclusivo. O algoritmo tendeu a buscar uma descrição meramente descritiva ou uma expansão textual volumosa do elemento "modo/meio", falhando em reconhecer que o detalhamento válido no modelo Enem pode ser aplicado de forma legítima

sobre qualquer um dos outros quatro elementos preexistentes. Essa rigidez mecânica da máquina pune o aluno que cumpre estritamente o edital, demonstrando que a inteligência artificial carece da flexibilidade interpretativa e do conhecimento técnico de banca necessários para validar construções textuais perfeitamente legítimas.

CONCLUSÃO

O presente estudo de caso descritivo e comparativo evidenciou que as ferramentas de Inteligência Artificial generativa baseadas em LLMs possuem um inegável potencial técnico de apoio e otimização ao ecossistema educacional contemporâneo. Elas atuam com extrema rapidez e precisão matemática em diagnósticos somativos superficiais, na identificação imediata de repetições lexicais brutas e no fornecimento de sugestões rápidas de substituição de conectivos gramaticais formais de superfície. Contudo, a IA demonstra limites analíticos severos e profundos ao lidar com as competências macroestruturais que exigem o exercício pleno da inteligência argumentativa abstrata, a atribuição de intencionalidade discursiva, a validação hermenêutica da legitimidade de repertórios socioculturais e a interpretação integrada de projetos de texto. Mais grave ainda, este estudo identificou casos de confabulação, em que uma das ferramentas atribuiu ao estudante repertórios socioculturais inexistentes no texto original — um risco pedagógico que ultrapassa a mera leniência de nota e atinge a própria fidedignidade do relatório de correção. A tendência sistemática das ferramentas de IA em conferir notas artificialmente infladas e emitir feedbacks vazios, genéricos ou factualmente fabricados representa um risco pedagógico grave, pois pode gerar no estudante uma falsa percepção de segurança intelectual e excelente desempenho, mascarando lacunas cognitivas profundas e prejudicando severamente a sua preparação real para o exame. É necessário reconhecer, no entanto, os limites do presente estudo. Trata-se de uma análise de caso único (N=1), sobre uma única redação e um único conjunto de respostas de cada ferramenta de IA em um momento específico no tempo; os resultados aqui discutidos, portanto, ilustram e evidenciam mecanismos de falha plausíveis,

mas não permitem generalizações estatísticas sobre o desempenho médio dessas ferramentas em corpora maiores. Embora o prompt utilizado esteja documentado no Apêndice D e as versões dos modelos empregados (ChatGPT, versão GPT-5.5, da OpenAI; e Gemini, versão 3.5 Flash) estejam agora identificadas, a data exata de acesso não foi formalmente registrada, o que ainda impede isolar eventuais efeitos de atualização dos modelos sobre os resultados. Estudos futuros poderão ampliar a robustez destes achados por meio de corpora maiores, de múltiplos avaliadores humanos independentes, de comparações entre diferentes versões e configurações de modelos, e de tratamento estatístico das divergências observadas. Dessa forma, conclui-se que o olhar sensível, analítico e experiente da mentoria de escrita personalizada consolida-se como um elemento metodológico particularmente relevante e, com base neste estudo de caso, não substituível pelas ferramentas de IA generativa avaliadas em seu estágio atual. O mentor humano demonstrou, neste caso, ser capaz de compreender a real intenção comunicativa do aluno, rastrear desvios lógicos complexos ocultos sob uma roupagem formal de palavras, acolher as demandas socioemocionais do estudante e intervir cirurgicamente por meio de uma avaliação verdadeiramente formativa e emancipatória.

REFERÊNCIAS

- BENDER, Emily M.; GEBRU, Timnit; MCMILLAN-MAJOR, Angelina; SHMITCHELL, Shmargaret. On the dangers of stochastic parrots: can language models be too big? In: PROCEEDINGS OF THE 2021 ACM CONFERENCE ON FAIRNESS, ACCOUNTABILITY, AND TRANSPARENCY (FAccT'21), p. 610-623, 2021.
- HADJI, Charles. Avaliação desmistificada. Porto Alegre: Artmed, 2001.
- HOFFMANN, Jussara. Avaliar para promover: as setas do caminho. Porto Alegre: Mediação, 2001.

INSTITUTO NACIONAL DE ESTUDOS E PESQUISAS EDUCACIONAIS ANÍSIO TEIXEIRA (INEP). A redação no Enem: manual de correção - competências de avaliação. Brasília: Inep, 2023.

INSTITUTO NACIONAL DE ESTUDOS E PESQUISAS EDUCACIONAIS ANÍSIO TEIXEIRA (INEP). Cartilha de Redação do Enem 2025. Brasília: Inep, 2025.

PACK, Austin; BARRETT, Alex; ESCALANTE, Juan. Large language models and automated essay scoring of English language learner writing: insights into validity and reliability. Computers and Education: Artificial Intelligence, v. 6, 100234, 2024.

RAWTE, Vipula; CHAKRABORTY, Swagata; PATHAK, Agnibh; SARKAR, Anubhav; TONMOY, S. M. Towhidul Islam; CHADHA, Aman; SHETH, Amit; DAS, Amitava. The troubling emergence of hallucination in large language models: an extensive definition, quantification, and prescriptive remediations. In: PROCEEDINGS OF THE 2023 CONFERENCE ON EMPIRICAL METHODS IN NATURAL LANGUAGE PROCESSING (EMNLP), Singapura, p. 2541-2573, 2023.

SHERMIS, Mark D.; BURSTEIN, Jill (Eds.). Handbook of automated essay evaluation: current applications and new directions. Nova York: Routledge, 2013.

APÊNDICE A – Redação Transcrita do Estudante

Caminhos para enfrentar a normalização dos jogos de azar feita pela publicidade no Brasil

Competências C1: 120 C2: 120 C3:160 C4: 120 C5: 200 Pontuação final: 720

Comentário final:O seu texto demonstra que você compreendeu o tema e soube estruturar a redação dentro dos moldes exigidos pelo ENEM. A sua proposta de intervenção também está completa, contendo todos os 5 elementos necessários de maneira bem desenvolvida. No entanto, sua nota foi penalizada principalmente pela C2. Você basicamente recorreu a "repertórios de bolso", frases decoradas que pouco reforçavam sua linha de raciocínio, servindo apenas como ornamentos do texto. Em C1, há muitos desvios ortográficos (falta de acentos, erros de grafia), problemas de concordância e falhas graves de pontuação (como separar sujeito do verbo com vírgula). Além disso, na C3 e na C4, você cometeu alguns tropeços ao tentar conectar ideias diferentes num mesmo período, o que gerou truncamentos e saltos argumentativos (como no final do primeiro parágrafo de desenvolvimento). Sugiro agendar uma sessão de Mentoria focada em revisão gramatical, na construção de períodos sintáticos mais claros e bem delimitados e na seleção de repertório. Você tem muito potencial, continue praticando!

"O mais escandaloso nos escândalos, é1 que nos habituamos a eles". Essa afirmação da filósofa2 francesa Simone de Beauvoir pode facilmente ser aplicada ao debate sobre a normalização da publicidade dos jogos de azar no Brasil, 30uma vez que mais escandaloso35 do que o avanço dessa realidade é o fato da sociedade3 compreender e29ssa situação com naturalidade. Esse cenário é fruto inegável de uma sistemática negligência do Estado 31que não promove leis de fiscalizações eficientes que barrem a intensa publicidade. Ademais4 é fundamental reconhecer como a mídia contribui ao não oferecer um olhar crítico5 ao público, juntamente com uma educação tecnicista.32

Em primeira analis6e é importante reconhecer como a mídia promove a normalização do marketing de jogos de azar como bets, cassinos virtuais, roletas,e demais jogos sejam vistos com7o hobbies 33não alertando sobre suas consequências. Esse panorama se aproxima do que defendeu o jornalista Caco Barcellos, n8a qual declarou que "a culpa não é de quem não sabe, mas de quem não informa"34, já que a ausência de comprometimento dos veículos de comunicação com 9a bem-estar 10à população brasileira11a, faz com que a normalização desses jogo12s acarretem em vícios 13a populaçã36o. Isso ocorr14e, porque os veículos de comunicação são patrocinados por grandes empresas e tendem a promover seus interesses econômicos ao não abordar como a flexibilização das leis pelo Estado , impulsiona16m empresas como os bancos, que lucram com a constante movimentação e linhas de crédito15o.

Adicionalmente, é fundamental perceber como as instituições de ensino promovem um panorama em que a falta de letramento digital, juntamente com a falta de suporte psicológico cr17ia um ambiente propício a uma geração de jovens que usam e divulgam jogos de azar. Esse cenário se aproxima do que defendeu o político Junqueira Gerra37, na qu18al dita 38" a escola é a única alavanca capaz de elevar ao povo 19ao nível da moral ", uma vez que o devido suporte e orientação não oferecida ao ma20is básico grau de ensino acarreta a criação de futuros apostadores e publicitários através 21as red22es sociais. Porém es23se contexto advém da ausência de comprometimento o sistema educacional c24om o desenvolvimento socioemocional dos alunos em detrimento de um currículo conteúdodista.

Portanto, faz2526-se imperativo o reconhecimento de que a normalização dos jogos de azar tem como origem a negligência do Estado. Assim, para solucionar tal realidade é necessário qu27e o Governo Federal, por meio de um Decreto, crie o Programa Nacional Contra a Divulgação de Jogos de Azar39, com a finalidade de vetar novas divulgações perante multa. Este Programa40 deve ainda, por meio de leis a serem votadas no Congresso, estipular uma regulamentação para que a mídia e instituições de ensino tenham o compromisso legal de promoverem conteúdos contra os jogos de azar e sua divulga28ção.

Comentários:Foram ainda registrados diversos desvios pontuais de norma culta (C1): pontuação (separação indevida de sujeito e verbo, ausência de vírgulas após conectivos e em expressões intercaladas, como após "Estado" e "hobbies"), acentuação e ortografia ("filósofa", "crítico", "análise", "publicitários", "através", "conteudista", "necessário"), regência e crase ("vícios à população", "elevar o povo", "comprometimento do sistema", "das redes sociais") e concordância verbal e nominal ("acarrete", "o bem-estar da população", "não oferecidos", "impulsiona").

#7 C3: Há um truncamento sintático que prejudica a coerência e o encadeamento lógico desta ideia. Ao iniciar com "como a mídia promove", a oração perde a conexão com o final "sejam vistos". Sugestão de reescrita: "como a mídia promove a normalização [...] fazendo com que sejam vistos como hobbies". #8 C4: Uso inadequado do pronome relativo. Como você está se referindo ao jornalista Caco Barcellos, o correto seria usar "o qual" ou simplesmente reescrever como "ao declarar que...". #15 C3: Houve um salto argumentativo neste trecho. Você estava falando sobre a mídia e, abruptamente, passou a falar sobre as leis do Estado e os lucros dos bancos em um mesmo período, o que deixou a argumentação confusa e prejudicou o foco do parágrafo. É importante desenvolver uma ideia de cada vez com clareza.#18 C4: Novo desvio de coesão referencial. Para retomar uma pessoa (o político), não se usa "na qual dita". O ideal seria "o qual afirma" ou "que dita".#25 C3: Falha semântica/argumentativa. A expressão "em detrimento de" significa "em prejuízo de". O que você quis dizer é que o sistema foca no currículo conteúdodista e abandona o lado socioemocional. O correto seria algo como "em prol de um currículo conteúdodista".#28 C5: Sua proposta de intervenção está completa. Você apresentou o Agente (Governo Federal), a Ação (criar o Programa Nacional...), o Modo/Meio (por meio de um Decreto), o Efeito (com a finalidade de vetar novas divulgações) e o

Detalhamento (Este Programa deve ainda... estipular uma regulamentação...).#30 C2: Embora seja legitimada (e até pertinente) essa citação é pouco produtiva: trata-se de um "repertório de bolso" (em suas próprias palavras, ela "pode ser facilmente aplicada" a qualquer problema naturalizado na sociedade brasileira). Evite esse tipo de referência, procure sempre mobilizar conceitos e autoridades que realmente tenham alguma relação com o tema discutido. #32 C3: Reelabore: segundo a frase a "mídia não oferece uma educação tecnicista" (aliás, nem seria dever dela ofertar qualquer tipo de educação). #34 C2: Novo caso de "repertório de bolso": evite usar citações decoradas que não sejam relacionadas realmente ao tema específico do qual você trata. #35 C4: Evite a repetição de termos (embora ela seja aceitável na frase, por se tratar de um uma citação, no resto do parágrafo se trata de descuido e mau planejamento do texto).#36 C3: Note como a frase está circular e repetitiva. Reelabore-a.#37 C1: Guerra Junqueiro, embora também tenha sido político, é muito mais conhecido como um poeta português. O erro na grafia do nome e o desconhecimento de sua mais importante área de atuação já prejudicam sua avaliação em C2. Porém, o mais problemático é, novamente, o recurso ao "repertório de bolso": você usa mais uma vez uma frase genérica, sem ligação com o tema proposto, apenas para enfeitar o texto com uma citação que não se liga realmente à discussão. Em lugar disso, seria mais eficiente recorrer a fatos ou estatísticas que trouxessem alguma informação sobre o assunto. #39 C1: Explique melhor essa proposta: por que seria necessário criar um "programa contra a divulgação dos jogos", em vez de simplesmente proibi-la? Ou aos jogos?

APÊNDICE B – Devolutiva e Justificativa Completa da Inteligência

Artificial 1 (ChatGPT, versão GPT-5.5, da OpenAI)

Competência 1 (160/200): a IA identificou boa variedade vocabular e estrutura sintática relativamente complexa, mas apontou desvios pontuais de pontuação, concordância e escolhas vocabulares, além de repetições como "normalização dos jogos de azar".Competência 2 (200/200): considerou o texto plenamente adequado ao tema e elogiou o repertório sociocultural mobilizado, listando como referências Simone de Beauvoir, Lúcia Santaella, Paulo Freire e o Programa Nacional Contra Publicidade de Jogos de Azar.Competência 3 (180/200) e Competência 4 (180/200): reconheceu boa progressão argumentativa e uso adequado de conectivos ("Em primeira análise", "Adicionalmente", "Portanto"), com ressalva a argumentos por vezes genéricos e repetições vocabulares.Competência 5 (180/200): avaliou a proposta de intervenção como completa, porém sugeriu maior detalhamento do modo de execução e da fiscalização. Nota final estimada pela IA 1: 900/1000 pontos.

APÊNDICE C – Devolutiva e Justificativa Completa da Inteligência

Artificial 2 (Gemini, versão 3.5 Flash)

Competência 1 (160/200): destacou domínio da norma culta e vocabulário rico, com poucos desvios gramaticais pontuais de grafia, concordância ou regência.Competência 2 (200/200): considerou os repertórios de Simone de Beauvoir, Caco Barcellos e Junqueira Freire extremamente produtivos, legítimos e bem articulados ao tema.Competência 3 (200/200) e Competência 4 (200/200): avaliou o projeto de texto como claro e estratégico, com progressão fluida entre os parágrafos e coesão interparágrafos e intraparágrafos exemplar.Competência 5 (200/200): considerou a proposta de intervenção completa, contemplando todos os elementos exigidos pela matriz do Enem (Agente, Ação, Meio, Efeito e Detalhamento). Nota final atribuída pela IA 2: 960/1000.

APÊNDICE D – Prompt Utilizado nas IAs Generativas

Reproduz-se a seguir, na íntegra, o comando (prompt) utilizado de forma idêntica para instruir a IA 1 (ChatGPT, versão GPT-5.5, da OpenAI) e a IA 2 (Gemini, versão 3.5 Flash), conforme descrito na seção Material e Métodos:"Atue como corretor oficial do Enem. Corrija a redação a seguir dando nota de 0 a 200 para cada uma das 5 competências e justifique."O mesmo

prompt foi seguido, em ambos os casos, pelo texto integral da redação transcrita (Apêndice A), sem instruções adicionais de refinamento além do comando acima.