

## APRESENTAÇÃO

LINGÜÍSTICA DE CORPUS E INTELIGÊNCIA ARTIFICIAL: INTERFACES COM LETRAS & LINGÜÍSTICA

---

***Corpus Linguistics and Artificial Intelligence:  
Interfaces with Modern Languages & Linguistics***

DOI: 10.14393/LL63-v41S-2026-00

Elisa Duarte Teixeira\*

Rozane Rodrigues Rebechi\*\*

Stella E. O. Tagnin\*\*\*

---

\* Doutora em Estudos Linguísticos e Literários em Inglês (USP). Professora de Tradução – Inglês do Departamento de Línguas Estrangeiras e Tradução da Universidade de Brasília. ORCID: 0000-0003-3472-3605. E-mail para contato: elisadut(AT)unb.br.

\*\* Doutora em Estudos Linguísticos e Literários em Inglês (USP). Professora de Tradução – Inglês do Departamento de Línguas Modernas da Universidade Federal do Rio Grande do Sul. ORCID: 0000-0002-1878-7548. E-mail para contato: rozane.rebechi(AT)ufrgs.br.

\*\*\* Livre-Docente em Estudos Linguísticos e Literários em Inglês (USP). Professora Sênior atuando em dois programas de pós-graduação na USP: Estudos Linguísticos e Literários em Inglês (PPGELLI) e Línguas Estrangeiras e Tradução (LETRA). ORCID: 0000-0002-5517-2710. E-mail para contato: seotagni(AT)usp.br.

**RESUMO:** Este texto apresenta uma coletânea de 14 artigos completos submetidos à chamada para publicação deste número especial da *Revista Letras & Letras*, dedicado à divulgação de trabalhos apresentados durante o evento conjunto ELC/EBRALC 2024 - XVI Encontro de Linguística de Corpus e XII Escola Brasileira de Linguística Computacional, realizado de 21 a 24 de outubro de 2024, na Universidade de Brasília. Os temas abordados vão dos estudos do léxico à apresentação de novos recursos para pesquisas com corpora, passando pela análise linguística baseada em corpus, em suas diversas aplicações, e por reflexões acerca do impacto da IA nos estudos com corpora e áreas conexas.

**PALAVRAS-CHAVE:** Linguística de Corpus; Processamento da Linguagem Natural; Linguística Computacional; Inteligência Artificial; Letras.

**ABSTRACT:** This introductory text presents a collection of 14 articles submitted to this special issue of *Revista Letras & Letras*, aimed at disseminating research presented at the 2024 ELC/EBRALC event—the 16th Corpus Linguistics Conference (ELC) and the 12th Brazilian School of Computational Linguistics (EBRALC)—held from October 21 to 24, 2024, at the University of Brasília. The topics covered range from lexical studies to the proposal of new resources for corpus-based research, corpus-based linguistic analyses in their various applications, and reflections on the impact of artificial intelligence on corpus linguistics and related fields.

**KEYWORDS:** Corpus Linguistics; Natural Language Processing; Computational Linguistics; Artificial Intelligence; Modern Languages.

---

A Linguística de Corpus (LC), campo disciplinar que estuda a linguagem por meio da análise de grandes quantidades de textos em formato eletrônico organizados na forma de corpora, avançou muito desde os primeiros trabalhos publicados, na década de 1960. Hoje, o número de áreas do conhecimento que se valem da LC em suas pesquisas, seja como abordagem, seja como metodologia, é cada vez maior e vai muito além das áreas de Letras e Linguística.

O Processamento de Linguagem Natural (PLN), ou Linguística Computacional, é uma área multidisciplinar que aplica a ciência da computação à análise e compreensão da linguagem humana. Esse campo teve um crescimento exponencial nas últimas décadas, à medida que o uso da tecnologia e, mais recentemente, da Inteligência Artificial Generativa, foram evoluindo e se fazendo presentes em praticamente todas as esferas da experiência humana.

Há mais de 25 anos, o evento conjunto ELC/EBRALC tem promovido o diálogo entre LC e PLN. O Encontro de Linguística de Corpus (ELC) teve sua primeira edição em 1999, na Universidade de São Paulo; a Escola Brasileira de Linguística Computacional (EBRALC) surgiu alguns anos depois, em 2007, à medida que participantes do ELC constataram uma necessidade premente de ampliar os conhecimentos computacionais e tecnológicos de pesquisadoras e pesquisadores do Brasil trabalhando com LC nas áreas de Humanas. Assim, o objetivo da EBRALC é oferecer oficinas práticas, ao passo que o ELC é um evento científico que tem por tradição não possuir sessões paralelas, para que todas as pessoas participantes possam assistir a todas as comunicações, proporcionando um convívio mais profícuo entre as várias áreas envolvidas – um convite ao diálogo, à reflexão e à colaboração.

O ELC/EBRALC 2024 – XVI Encontro de Linguística de Corpus e XII Escola Brasileira de Linguística Computacional –, realizado de 21 a 24 de outubro de 2024 na Universidade de Brasília, teve como tema “Linguística de Corpus e Inteligência Artificial: interfaces com Letras e Linguística” e foi a primeira edição do evento depois do hiato ocasionado pela pandemia de COVID-19, que prejudicou a periodicidade dos encontros presenciais. Organizado por docentes do Instituto de Letras da UnB membros dos grupos de pesquisa TermiTraDiCo (Terminologia e Tradução Direcionadas por Corpus) e CoMPLETT (Corpus Multilíngue para Pesquisas em Línguas Estrangeiras, Tradução e Terminologia), além de colegas de várias

outras universidades, o evento internacional contou com o apoio do Programa de Pós-Graduação em Estudos da Tradução do Instituto de Letras da UnB (POSTRAD), da Associação Brasileira de Linguística de Corpus (ABraLC) e da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES). A proposta era fomentar discussões sobre os impactos da Inteligência Artificial na formação acadêmica e na pesquisa nas áreas de Letras e Linguística, com foco no papel que a LC tem ocupado e ainda pode ocupar no estabelecimento desse diálogo transdisciplinar.

A programação contou com grandes nomes das áreas envolvidas. A EBRALC teve início com a palestra de abertura intitulada “O futuro da Linguística Computacional com as novas ferramentas de Inteligência Artificial”, proferida pelo Prof. Dr. Osvaldo Novais de Oliveira Jr., Professor Titular do Departamento de Física da Universidade de São Paulo e pesquisador-fundador do Núcleo Interinstitucional de Linguística Computacional (NILC-USP).

Na sequência, oficinas paralelas de 2 ou 4 horas versaram sobre temas diversos, com diferentes níveis de aprofundamento, para atender a todas as pessoas participantes. O próprio Prof. Osvaldo ministrou a oficina “Escrita científica em inglês com o uso de ferramentas inteligentes”, de 2 horas. Simultaneamente, Leo Selivan, um de nossos convidados internacionais, autor de livros e capítulos sobre o uso da LC no ensino do léxico da língua inglesa, ministrou a primeira edição de sua oficina “*Building Activities to Teach L2 Vocabulary with Corpus*”, oferecida também no segundo dia. Na parte da tarde, o Prof. Dr. Jean-Claude Miroir, da Tradução – Francês da UnB, ministrou “Processamento de Linguagem Natural Multilíngue com SpaCy e AntConc (v. 4)”, seguida pela oficina oferecida pela convidada Prof. Dra. Elaine Trindade, vice-coordenadora e professora do curso Letras-Tradutor da PUC-SP, intitulada “Lancsbox: uma ferramenta para explorar corpora com precisão e interatividade”. Luciana Latarini Ginezi, tradutora e intérprete forense com larga experiência na área e doutorado focado em Tradução e Terminologia, ministrou paralelamente sua oficina de 4 horas intitulada “Introdução às tecnologias de apoio à interpretação”.

No segundo dia, a pesquisadora do Centro de Inteligência Artificial (C4AI) da Universidade de São Paulo, Cláudia Freitas, autora de livros, artigos e capítulos sobre Linguística Computacional e membro do coletivo Brasileiras em PLN, ministrou a oficina

“Introdução à Linguística Computacional”, com duração de 4 horas. Na outra sala, o Prof. Dr. Cláudio Corrêa e Castro Gonçalves, docente do Departamento de Línguas Estrangeiras e Tradução da UnB, juntamente com a convidada Anna Beatriz Dimas Furtado, egressa da UnB com Mestrado em Tecnologia para Tradução e Interpretação pelo programa European Masters in Technology for Translation and Interpreting (EM TTI) e, à época, assistente de pesquisa em Ciência de Dados Linguísticos e Linguística de Corpus para Tradução da Universidade de Galway, na Irlanda (onde atualmente é doutoranda), ministraram a oficina de 4 horas “Segundos passos em Python: organizando a produção de programas”.

Na parte da tarde, o Prof. Dr. Leonel Figueiredo de Alencar, titular da Universidade Federal do Ceará e membro da organização *Universal Dependencies* do GitHub, dedicada à construção de corpora e tradução automática das línguas indígenas cadiuéu e nheengatu, ministrou a oficina de 4 horas “Construindo um analisador morfossintático no modelo *Universal Dependencies* (UD)”, dividida em duas sessões: uma introdutória, para quem tivesse interesse em aprender os conceitos básicos da UD, e outra voltada à aplicação desses modelos às línguas indígenas brasileiras. Paralelamente à parte mais avançada desta oficina, tivemos uma de 2 horas intitulada “*We need to talk about language technologies!*”, ministrada pela convidada internacional Profa. Dra. Lynne Bowker, professora titular e coordenadora do programa de pós-graduação *Research in Translation, Technologies and Society*, da Université Laval, no Canadá.

Bowker, que é Fellow of the Royal Society of Canada, além de autora, educadora e divulgadora científica de grande destaque nas áreas de LC, Tradução, Terminologia e divulgação científica, também ficou a cargo da palestra de abertura do ELC, ocorrida no dia 23/10, intitulada “*From Corpus Linguistics to AI: The enduring importance of having the right data*” (Da Linguística de Corpus à Inteligência Artificial: a importância constante de se trabalhar com dados confiáveis).

Nas sessões de comunicações orais do evento que se seguiram foram feitas 22 apresentações de 20 minutos, organizadas em 5 sessões, com espaço para debate ao final de cada uma. Também houve uma sessão de pôsteres na tarde do primeiro dia, com 19 trabalhos apresentados, além de 8 vídeos curtos pré-gravados, de 5 minutos cada, exibidos continuamente durante a sessão de pôsteres, na tela da sala principal do evento. O teor

desses trabalhos e a programação completa do evento podem ser apreciados nos Anais do ELC/EBRALC 2024, disponível em: <https://www.elc-ebralc.net.br/anais>.

Neste volume, reunimos 14 artigos completos de alguns dos trabalhos apresentados durante o ELC e submetidos à chamada para publicação de número especial da *Revista Letras & Letras*, cuja Diretoria aproveitamos para agradecer, em especial à pessoa do Prof. Dr. Igor Antônio Lourenço da Silva, pela acolhida e por nos acompanhar ao longo do trabalhoso processo de preparação da coletânea. Os temas abordados vão dos estudos do léxico à apresentação de novos recursos para pesquisas com corpora, passando pela análise linguística baseada em corpus, em suas diversas aplicações, e por reflexões acerca do impacto da IA nos estudos com corpora e áreas conexas.

Em “Criação de um glossário bilíngue libras-português: uma iniciativa para promover acessibilidade comunicacional”, Ronnie Fagundes de Brito e Gabriela Wick-Pedro apresentam as etapas iniciais do desenvolvimento de um glossário em português do Brasil e Língua Brasileira de Sinais (Libras), concebido para ampliar o acesso da comunidade surda à informação científica e tecnológica. O artigo descreve a metodologia adotada, que combina técnicas de PLN para a extração automática de termos com análise semântica e validação realizada por especialistas ouvintes e surdos. Como resultado, foram identificados, categorizados e validados 130 termos científicos, organizados em *clusters* temáticos e estruturados em uma base de dados terminológica que contempla suas relações semânticas. Ao evidenciar o potencial das tecnologias linguísticas para a construção de recursos terminológicos acessíveis, o trabalho contribui para a promoção da inclusão e da democratização do conhecimento científico.

O artigo “Glossários médicos multilíngues para CAT tools: desenvolvimento baseado na CID-11 da OMS”, de Jean-Claude Miroir, apresenta uma metodologia para transformar os dados oficiais da Classificação Internacional de Doenças (CID-11), da OMS, em glossários multilíngues compatíveis com ferramentas de tradução assistida por computador (CAT tools). Partindo da constatação de que a tradução médica demanda elevado grau de consistência terminológica e carece de recursos padronizados e interoperáveis, o autor descreve o processo de compilação de 36.049 pares terminológicos em inglês, português, francês e espanhol, estruturados nos formatos TBX e XLSX, para desenvolver o controle de

qualidade em traduções especializadas. O artigo também discute o potencial desses dados para exploração linguística por meio de concordanciadores e analisa o desempenho de ferramentas de PLN, como SpaCy e Stanza, enfatizando as limitações dos modelos generalistas na classificação de nomes próprios.

Amanda Letícia Valadares dos Santos e Flávia de Oliveira Maia-Pires, no artigo “Revisão e ampliação de árvores de domínio a partir da análise de *corpus*”, buscam demonstrar como a LC pode ser uma abordagem metodológica aliada na constituição, revisão e atualização de árvores de domínio terminológicas, por sua capacidade já demonstrada na literatura de auxiliar na identificação de candidatos a termo. As autoras discutem a importância das árvores de domínio como forma de arquitetar as informações de uma área de especialidade, auxiliando na identificação de conceitos e variações terminológicas e, especialmente, na compreensão das relações entre conceitos conexos. Partindo de um diagrama inicialmente elaborado pelas autoras a partir da leitura e análise manual da *Lei Geral de Proteção de Dados*, o trabalho descreve a exploração de um corpus de estudo sobre o mesmo tema utilizando as ferramentas de LC do Sketch Engine, em especial o Word Sketch. As autoras demonstram como a estratégia resultou na identificação de 12 novos termos essenciais que não constavam do diagrama inicial, promovendo um aprofundamento conceitual e um mapeamento mais preciso das correlações terminológicas da área em foco.

No artigo “Classificadores semânticos na identificação de equivalentes interlinguais para expressões idiomáticas: descascando esse abacaxi”, Julia de Souza Batista, Isabela Moreira de Oliveira, Stella Esther Ortweiler Tagnin, Elisa Duarte Teixeira e Rozane Rodrigues Rebechi apresentam uma proposta metodológica para o desenvolvimento de um sistema de descritores semânticos voltado à identificação semiautomática de equivalentes interlinguais entre expressões idiomáticas relacionadas à alimentação. Inserido em um projeto de construção de um dicionário onomasiológico multilíngue, o estudo descreve a elaboração e a testagem inicial desse sistema por meio da atribuição de descritores a expressões idiomáticas em português e inglês. Os resultados revelam baixo índice de consenso entre as(os) anotadora(e)s humanos, evidenciando a complexidade da representação semântica desse tipo de unidade fraseológica e os desafios conceituais e metodológicos que ainda

precisam ser enfrentados para viabilizar a identificação automática de expressões idiomáticas equivalentes entre línguas.

Em “A criação do MACC: Machado de Assis Catálogo & Corpus”, Ursula Puello Sydio descreve as etapas de planejamento e compilação desse recurso digital que oferece um catálogo abrangente e um corpus bilíngue das obras de Machado de Assis traduzidas para o inglês, recurso inexistente, até então. O Catálogo reúne 40 títulos do autor, entre romances, contos e antologias, organizados em uma linha do tempo interativa que permite buscar por títulos em português ou inglês. O Corpus eletrônico paralelo bilíngue, por sua vez, é constituído de 6 romances de Machado alinhados com suas 11 traduções para o inglês e mais 90 contos alinhados a 373 traduções. O artigo também descreve as etapas de desenvolvimento do *website* de consulta a esse valioso recurso, além de sugerir como o MACC pode ser usado por pesquisadores, tradutores e demais interessados.

Em “CALIEMT: corpus de aprendizes da língua inglesa do ensino médio técnico”, Lucas Renato dos Santos Murta e Shirlene Bemfica de Oliveira descrevem a composição de um corpus de produções orais e escritas de estudantes de inglês do ensino médio técnico de um Instituto Federal. Integrado a um projeto de ensino, pesquisa e extensão, o corpus reúne dados produzidos em atividades colaborativas voltadas ao desenvolvimento da proficiência linguística, do pensamento crítico e da autonomia dos aprendizes. O artigo descreve as etapas de organização do corpus, os critérios de seleção dos metadados e apresenta uma análise inicial inspirada na Análise Multidimensional. Ao disponibilizar um conjunto representativo de produções de aprendizes em diferentes níveis de proficiência, o CALIEMT amplia as possibilidades de investigação sobre a aprendizagem de inglês e a constituição de corpora de aprendizes no contexto brasileiro.

A construção de um corpus de notícias satíricas em português brasileiro, desenvolvido a partir de publicações do site *Sensacionalista*, é apresentada por Gabriela Wick-Pedro na contribuição “SatiriCorpus.Br: um corpus de notícias satíricas para o português do Brasil”. Partindo da distinção entre sátira (que tem como objetivo provocar a reflexão e o riso) e *fake news* (disseminadas com a intenção de enganar), a autora discute as características desse gênero discursivo, que utiliza humor, ironia e exagero para promover crítica social. O artigo descreve a compilação do corpus e destaca seu potencial para a

investigação de padrões lexicais, sintáticos e semânticos característicos da linguagem satírica, oferecendo um recurso inédito para pesquisas em LC e PLN voltadas à detecção automática desse tipo de conteúdo.

Deise Prina Dutra, Danilo Duarte Costa, Carolina Godoi de Faria Marques, Gustavo Leal Teixeira, Ana Eliza Pereira Bocorny e Jhonatan Henrique Lopes Alves apresentam o CROPS (*Corpus of Agrarian Sciences Postgraduate Studies*) no artigo intitulado “Introducing CROPS: a Specialized Corpus of English-language Agrarian Sciences Postgraduate Theses and Dissertations from Brazilian Universities”, um corpus especializado composto por dissertações e teses escritas em inglês por estudantes de programas de pós-graduação brasileira(o)s da área de Ciências Agrárias. A contribuição descreve as etapas de compilação e organização do CROPS, que reúne aproximadamente 2,4 milhões de palavras distribuídas em diferentes subáreas das Ciências Agrárias, constituindo um recurso representativo da escrita acadêmica produzida nesse campo. Conforme mostram as pessoas autoras, o recurso oferece novas possibilidades para pesquisas linguísticas e estudos sobre o inglês acadêmico.

Patrícia Helena Freitag, em “O XBENCH como um recurso para garantir a convencionalidade na tradução de colocações”, demonstra como ferramentas de controle de qualidade podem ser empregadas para identificar, de forma automatizada, estruturas não convencionais em traduções especializadas. Com base na tradução de artigos de blogs de *coworking*, a autora descreve a elaboração de uma *checklist*, construída a partir de padrões de colocações inadequadas identificadas em pesquisas anteriores, para o uso no software de controle de qualidade. O artigo apresenta exemplos de regras formuladas com expressões regulares, capazes de detectar traduções pouco convencionais e auxiliar o trabalho de revisão. Além de evidenciar o potencial do Xbench para promover maior naturalidade e consistência nas traduções, a autora destaca que a metodologia proposta pode ser adaptada para diferentes gêneros textuais e domínios especializados.

No artigo “Quão confiáveis são as ferramentas de IA para a tradução de receitas culinárias? Algumas surpresas”, Stella Tagnin e Rozane Rebechi comparam o desempenho de ferramentas de tradução automática neural, como o Google Tradutor e o DeepL, com chatbots de IA generativa e as traduções publicadas de uma receita culinária italiana do século XIX para o português brasileiro e o inglês. As autoras usam corpora como o COCA e o

Corpus do Português, de Mark Davies, além de uma parte do corpus de culinária brasileira que serviu de base para a tese de Rebechi (2015), para atestar a convencionalidade e a recorrência das soluções oferecidas em textos originalmente produzidos por falantes nativos. Os resultados mostram que, apesar de as ferramentas computadorizadas sugerirem soluções surpreendentes em alguns pontos da receita, a tradução ainda não pode prescindir da pós-edição humana, especialmente nessa área específica do conhecimento, com grande carga cultural.

Roana Rodrigues, Paula Christina Figueira Cardoso e Jackson Wilke da Cruz Souza, em “Avaliação de LLMs na identificação de sinalizadores discursivos em textos jornalísticos anotados com RST”, investigam o potencial de Grandes Modelos de Linguagem na identificação automática de sinalizadores discursivos em textos jornalísticos do português brasileiro anotados segundo a *Rhetorical Structure Theory* (RST). Comparando as anotações produzidas pelos modelos GPT-4 e LLaMA-4 com aquelas realizadas por especialistas humanos, os autores observam uma convergência na identificação dos sinalizadores mais prototípicos das relações retóricas analisadas. Enquanto o GPT-4 apresentou anotações mais detalhadas, incluindo elementos nem sempre diretamente relacionados às relações discursivas, o LLaMA-4 produziu resultados mais concisos e, em alguns casos, sugeriu novos sinalizadores ainda não catalogados. Os resultados indicam que, embora os LLMs não substituam a anotação humana, eles podem atuar como ferramentas de apoio, tornando o processo de anotação mais eficiente e ampliando as possibilidades de investigação em estudos discursivos e linguística computacional.

Em “*Key-Feature Analysis for Dummies: A Comprehensive Step-By-Step Guide for Novice Researchers*”, Carolina Bohórquez Grondona e Luciana Aguiar de Oliveira apresentam um guia prático para a análise de Key-Features, ou características linguísticas (geralmente categorias léxico-gramaticais) cuja frequência é estatisticamente maior ou menor em um corpus, quando comparada a um corpus de referência. Com essa pesquisa, as autoras almejam tornar essa metodologia mais acessível a pesquisadores iniciantes em LC. Além de explicar, de forma detalhada, os fundamentos da técnica, os cálculos estatísticos e as ferramentas computacionais empregadas, as autoras demonstram sua aplicação por meio da comparação entre introduções de artigos científicos e de teses escritas em inglês. A análise

ênfatiza diferenças no uso de padrões léxico-gramaticais entre os dois gêneros, tais como a maior frequência de construções em que um advérbio ou outro elemento é inserido entre um verbo auxiliar e o verbo principal – split auxiliaries – e de verbos no infinitivo em artigos e de pronomes demonstrativos em teses, revelando características discursivas distintas.

Em “*Nominalizations as head nouns or modifiers: a corpus-based study of discussion sections in Forestry research articles*”, Shirlene Bemfica de Oliveira, Deise Prina Dutra e Jasper Vilan Braga tratam da escrita de artigos científicos, com foco específico nos sintagmas nominais nominalizados presentes na seção de discussão de artigos de pesquisa em Engenharia Florestal. Com vistas a informar pesquisadores da área e contribuir para a divulgação internacional da produção científica brasileira, o trabalho identifica os sufixos nominalizadores mais comuns, extraídos de um corpus constituído de 90 artigos de alto impacto da área de Engenharia Florestal, e discute como esses sintagmas se comportam na referida seção dos artigos, em termos de estrutura e função discursiva. Os resultados revelam uma predominância de sintagmas nominais nominalizados incorporados, o que aumenta a impessoalidade do texto, conforme observado anteriormente por outras pesquisas. É interessante observar como a alta frequência desse padrão, que promove um adensamento da informação e um maior nível de abstração e complexidade gramatical, pode ser percebida como um sinal de qualidade textual, conforme apontam as pessoas autoras, potencialmente impactando na aceitação de trabalhos para publicação.

Em “Desafios da Linguística de Corpus impostos pela Inteligência Artificial: rediscutindo alguns conceitos”, Jackson Wilke da Cruz Souza promove uma reflexão sobre os impactos da inteligência artificial, em especial dos Grandes Modelos de Linguagem, para conceitos centrais da Linguística de Corpus, como autenticidade, naturalidade e representatividade. A partir de uma revisão da literatura recente, o autor discute os desafios metodológicos decorrentes da crescente presença de textos gerados por IA, questionando seus efeitos sobre a construção e a confiabilidade de corpora. O artigo também aborda as implicações dessa transformação para noções como autoria, identidade e produção científica, propondo caminhos para que a LC possa incorporar, de forma crítica, os avanços da IA sem comprometer seus fundamentos teóricos e metodológicos.

Esperamos que os artigos apresentados neste número especial com algumas das contribuições apresentadas durante o ELC/EBRALC 2024 na UnB possam contribuir para o fortalecimento da LC e do PLN no Brasil e, principalmente, fomentar novas colaborações e avanços na intersecção dessas duas áreas de especialidade.

Aproveitamos para agradecer imensamente a todas e todos os membros da Comissão Científica e, especialmente, a colegas pesquisadoras e pesquisadores, membros ou não da Comissão, que contribuíram com seus pareceres ao longo de todo o processo de seleção de trabalhos para o evento e para este número especial – não se faz divulgação de ciência sem sua prestimosa contribuição.

Brasília, julho de 2026.

Elisa Duarte Teixeira  
Rozane Rodrigues Rebechi  
Stella E. O. Tagnin

Recebido em: 01.07.2026

Aprovado em: 27.07.2026