

Transformer-Based Multimodal Human Attribute Recognition and Identity Matching Under Appearance Variations

Wijay Sangh

Department of Computer Science, George Mason University, Fairfax, VA, USA.
contactvijay@gmu.edu

Abstract

The reliable recognition of human attributes and the matching of identities across non-overlapping camera views remain among the most challenging tasks in large-scale visual surveillance and intelligent infrastructure systems. Appearance variations induced by illumination changes, pose deformations, occlusions, and especially clothing alterations fundamentally undermine the effectiveness of conventional unimodal person re-identification and attribute analysis pipelines. This paper presents a comprehensive systems-level examination of transformer-based multimodal architectures that integrate heterogeneous sensory streams—visible, infrared, depth, and gait modalities—to achieve robust identity matching and attribute inference under severe appearance shifts. We analyze the structural trade-offs inherent in early, intermediate, and late fusion strategies within cross-modal transformer backbones, highlighting the role of decoupled attention mechanisms, modality-specific tokenization, and hierarchical feature propagation in balancing representational richness with computational efficiency. Beyond architectural considerations, the paper interrogates the governance, fairness, and sustainability implications of deploying such systems at scale, addressing demographic bias amplification, privacy erosion, and energy consumption. We further discuss deployment infrastructures spanning edge-cloud continuums, federated learning paradigms for privacy-preserving model updates, and regulatory constraints imposed by data protection frameworks such as the GDPR. By synthesizing perspectives from computer vision, systems engineering, and socio-technical policy studies, this work articulates a forward-looking research agenda that prioritizes not only algorithmic accuracy but also systemic resilience, ethical accountability, and long-term operational viability in real-world identity inference ecosystems.

Keywords

multimodal transformers, person re-identification, attribute recognition, appearance variation, cross-modal attention, system governance, biometric fairness, edge deployment.

1. Introduction

The proliferation of networked sensing infrastructures across transportation hubs, commercial spaces, and smart urban environments has generated an unprecedented demand for automated systems that can recognize human attributes and link person identities across disjoint camera views. Applications ranging from forensic video analysis to adaptive retail analytics and contactless access control depend on the ability to perform robust person re-identification and fine-grained attribute inference under widely varying imaging conditions. However, the intrinsic variability of human appearance—driven by changes in clothing, illumination, viewpoint, and partial occlusions—poses a persistent obstacle that has not been fully resolved

by unimodal deep learning models trained on standard RGB benchmarks. The emergence of transformer architectures, originally conceived for natural language processing [1] and subsequently adapted to vision tasks [2], has opened new pathways for designing highly expressive models capable of capturing long-range spatial dependencies and cross-modal interactions that are essential for appearance-robust person analysis.

From a systems engineering standpoint, the challenge extends beyond the mere optimization of recognition accuracy on curated test sets. Real-world identity inference infrastructures must simultaneously address model scalability, real-time throughput, energy efficiency, demographic fairness, and compliance with evolving data protection regulations. Unimodal person re-identification surveys [3] have thoroughly documented the limitations of early convolutional network designs, while pedestrian attribute recognition methods [4] have demonstrated that auxiliary semantic cues can improve matching performance. Yet the systematic integration of multiple sensing modalities within a unified transformer framework remains underexplored from an architectural and governance perspective. Multimodal approaches that fuse RGB, infrared, depth, and even gait sequences offer a principled means of mitigating appearance variability because each modality provides complementary invariance properties: thermal signatures are relatively insensitive to illumination and clothing texture, depth encodes geometric structure that persists across outfit changes, and gait patterns capture idiosyncratic motion dynamics that are difficult to conceal.

This paper adopts a holistic, infrastructure-aware lens to examine the design, deployment, and oversight of transformer-based multimodal human attribute recognition and identity matching systems. We avoid narrow algorithmic benchmarking and instead foreground the structural trade-offs that emerge when cross-modal attention mechanisms are scaled to operate across heterogeneous sensor fleets, when decoupled feature pathways are introduced to separate identity-relevant signals from transient appearance factors, and when these systems are embedded within larger sociotechnical assemblages that include edge processing nodes, cloud orchestration layers, and human-in-the-loop governance protocols. By articulating the interplay between representational capacity, computational resource allocation, and regulatory obligations, we aim to provide a comprehensive framework for researchers, system integrators, and policy analysts who must navigate the complex terrain of next-generation identity-aware infrastructures.

2. Related Work and System Context

The evolution of person re-identification has progressed from handcrafted feature ensembles to deep metric learning, and more recently to transformer-based architectures that leverage self-attention to model spatial part correlations. Early cross-modality approaches demonstrated that aligning RGB and infrared representations through shared feature subspaces could significantly improve matching in low-light environments [5], while Transformer-based re-identification models such as TransReID [6] introduced pure attention backbones that outperform their convolutional counterparts by capturing holistic body-structure relationships. Complementary work on gait-based recognition [7] revealed that silhouette sequences convey identity signatures that remain stable across weeks, providing a modality inherently resistant to clothing changes. Video-based person re-identification further exploits temporal attention mechanisms to aggregate evidence across frames, thereby reducing the impact of momentary occlusions and pose variations [8]. These individual strands underscore a crucial insight: no single modality or temporal granularity suffices to handle the full spectrum of appearance perturbations encountered in operational settings.

Despite these advances, the vast majority of systems studied in the literature are evaluated under closed-world assumptions with fixed camera networks, controlled lighting, and artificially clean identity labels. The systemic challenges of unconstrained deployment—sensor heterogeneity, asynchronous data streams, bandwidth constraints, and adversarial environmental conditions—remain peripheral to most algorithm-centric investigations. Moreover, the governance implications of fusing multiple biometric modalities are rarely addressed within the same analytical frame. The integration of cross-modal transformers in re-identification, while promising, introduces new attack surfaces related to modality spoofing and sensor tampering, demanding a reassessment of trust boundaries at the architectural level.

3. Multimodal Transformer Architecture and Fusion Strategies

A central architectural decision in transformer-based multimodal human analysis concerns the point at which information from diverse modalities is combined. Early fusion, which concatenates tokenized observations from RGB, infrared, and depth channels at the input embedding layer, allows the self-attention mechanism to freely compute inter-modal correlations from the very first layer. This approach maximizes representational flexibility but incurs a quadratic growth in sequence length, imposing severe memory and latency penalties that challenge edge deployment on resource-constrained devices. Late fusion, by contrast, processes each modality through independent transformer encoders and merges their outputs only at the final classification or embedding stage. While computationally more tractable, late fusion can fail to capture the nuanced cross-modal interactions that are critical for disambiguating identities when a single modality is compromised—for example, when a person dons a heavy coat that masks body shape in infrared while partially occluding RGB appearance.

Intermediate hierarchical fusion strategies offer a compelling middle ground. In such designs, modality-specific shallow transformers extract local features that are progressively aggregated through cross-attention blocks at multiple depths. This architecture creates a structured gradient of multimodal integration, preserving modality-specific invariances in early layers while allowing higher semantic concepts—such as joint representations of body geometry, thermal silhouette, and clothing texture—to be refined through shared attention heads. Recent work on part-aware transformers for visible-infrared person re-identification [10] exemplifies how structured attention can be directed toward body part configurations that remain consistent even when clothing changes, by learning part-level correspondences across modalities. Similarly, disentangled representation learning techniques [9] have shown that factorizing identity-relevant and appearance-relevant latent variables stabilizes matching performance under drastic visual transformations.

A critical systems-level consideration is the dynamic adaptation of fusion depth based on available computational resources and sensor quality. In a distributed surveillance network, a camera node with an embedded GPU may perform late fusion locally and transmit compact embeddings to a cloud aggregator that conducts deeper cross-modal re-ranking. The same architecture can be re-parameterized at inference time without retraining, provided that the underlying tokenization and attention modules are designed in a modular fashion. This plasticity is essential for sustaining performance across heterogeneous hardware fleets and for graceful degradation when certain modalities become intermittently unavailable due to sensor failure or privacy-preserving obfuscation. Such adaptability is especially relevant given the results from part-based visible-infrared re-identification studies [10], where part-level cross-attention proved robust to missing modality segments.

Decoupled attention mechanisms represent another architectural innovation that addresses clothing-induced appearance variations directly. By partitioning the feature space into a modality-specific stream that captures transient appearance cues and a modality-shared stream that encodes persistent identity attributes, the model can learn to suppress clothing-related activations during identity matching while still leveraging those cues for attribute prediction tasks. The required decoupling is achieved through dedicated attention pathways that are regularized to minimize mutual information between the two streams. This approach does not require explicit clothing labels during training, instead relying on contrastive objectives that encourage identity invariance across clothing conditions. When integrated into a multimodal transformer, the decoupled streams can be cross-attended at carefully chosen fusion points, preventing the propagation of appearance noise from one modality to another while preserving the complementary information that each modality contributes.

4. Attribute Recognition and Identity Matching Under Appearance Variations

Human attribute recognition and identity matching, though often treated as separate tasks, are deeply intertwined in transformer-based multimodal systems. Attributes such as gender presentation, age group, carried accessories, and clothing category serve both as retrieval filtering criteria and as auxiliary supervision signals that regularize the shared representation space. Under appearance variations, however, the reliability of attribute predictions can degrade sharply: a change of jacket alters clothing color attributes, while a hat can occlude hairstyle cues. The same transformer backbone must therefore learn to predict attributes that are stable across encounters—body height from depth, gait cadence from silhouette sequences—while also inferring transient attributes that may be useful for short-term re-identification but become misleading over longer intervals.

A transformative systems design principle for handling clothing changes is the explicit temporal modeling of attribute stability. Attributes derived from skeletal geometry and motion patterns exhibit low temporal variance and can be treated as long-term biometric anchors, whereas textile attributes have high variance and should be discounted or actively decorrelated from the identity embedding. Incorporating recurrent temporal aggregation modules on top of the multimodal transformer allows the system to accumulate evidence for stable attributes across a tracklet while gradually forgetting transient ones. The graph of attributes can also be structured hierarchically, with causal dependencies encoded through masked attention patterns that prevent the model from learning spurious correlations between clothing color and identity when trained on datasets with limited subject populations.

In this context, decoupling feature-driven and multimodal fusion attention for clothing-changing person re-identification has been framed as a specialized architectural intervention [11]. The approach separates the learning of modality-specific appearance transformations from identity-consistent fusion, ensuring that clothing variations are modeled in a contained feature subspace that does not contaminate the shared identity representation. Such decoupling provides a clean interface for incorporating new modalities: an additional sensor stream can be attached to the fusion attention block without retraining the appearance decoupler, simplifying system evolution. Beyond the algorithmic level, the decoupled design facilitates auditability, because the contributions of appearance versus identity pathways can be inspected separately to detect whether demographic attributes are inadvertently entangled with identity decisions—an important capability for fairness verification.

Long-term cloth-changing re-identification benchmarks [13] have spurred the development of generation-based augmentation strategies, wherein generative models synthesize plausible

clothing alterations to expand training distributions [14]. While effective in simulation, these data augmentation pipelines introduce system-level risks when transferred to operational contexts where the distribution of clothing types differs markedly from the training corpus, potentially amplifying biases against cultural attire. Attribute-driven feature disentanglement for video-based re-identification [12] offers a complementary pathway by forcing the network to explicitly separate attribute manifolds from identity manifolds, thereby reducing reliance on synthetic data. From a deployment perspective, a hybrid strategy that combines generative augmentation with disentangled representation learning and the decoupled fusion attention [11] constitutes a robust foundation for handling appearance variations while retaining the flexibility to incorporate new attributes and modalities as sensing capabilities evolve.

Privacy-preserving measures [15] add another layer of complexity to attribute recognition and identity matching under appearance variations. Techniques such as person-specific adversarial obfuscation that prevent unauthorized re-identification must operate without destroying the attribute cues that legitimate applications require—for instance, enabling customer counting in retail while preventing individual tracking. Transformer architectures can accommodate this dual requirement by partitioning the embedding space into public and private subspaces, with public attributes and private identity features processed through separate attention heads that are cryptographically isolated at deployment. The decoupled attention designs discussed earlier, including the feature-driven decoupling for clothing-changing re-identification [11], provide an architectural scaffold for such privacy-aware functional separation, illustrating how system-level requirements can directly shape neural network topology.

5. Infrastructure, Scalability, and Deployment Considerations

The transition from laboratory prototypes to real-world multimodal identity inference systems demands an infrastructure that reconciles the computational intensity of transformer architectures with the operational constraints of edge computing nodes, bandwidth-limited networks, and energy budgets. Edge computing paradigms [16] advocate for distributing inference workloads across a hierarchy of processing tiers, from low-power sensor-embedded microcontrollers that perform early-exit tokenization to gateway servers that execute full-depth cross-modal attention. The architectural trade-offs discussed in Section 3 directly impact this tiered deployment: late fusion configurations can be partially offloaded to cloud aggregators with minimal edge-side computation, while early fusion models necessitate more powerful edge accelerators but reduce data transmission volumes by fusing features close to the sensor origin.

Model compression techniques [17] such as knowledge distillation, quantization, and pruning are indispensable for shrinking multimodal transformers to fit edge constraints. However, naively compressing cross-modal attention blocks can disproportionately degrade the fusion quality for modalities that are naturally lower resolution or noisier, leading to uneven robustness across sensing conditions. This requires modality-aware compression schedules that allocate higher precision to infrared-depth cross-attention layers, for example, while aggressively quantizing high-resolution RGB early layers that are relatively redundant. System designers must also consider the sustainability implications of large-scale model serving: the carbon footprint of continuously operating thousands of inference nodes can be substantial, and recent calls for Green AI [23] emphasize the need to balance accuracy gains against energy proportionality. Selecting fusion depths dynamically based on query difficulty—spending computation only when appearance ambiguity is high—offers a promising route toward energy-proportional multimodal inference.

Federated learning [18] provides a privacy-compliant framework for updating person re-identification and attribute recognition models across distributed camera networks without centralizing raw biometric data. In such a setup, each site trains a local model on its sensor streams and shares only encrypted gradient updates with an aggregation server. The multimodal nature of the task introduces new challenges: different sites may possess heterogeneous subsets of modalities, making naive parameter averaging suboptimal. Cross-modal knowledge distillation within the federated curriculum can align representation spaces across sites, while decoupled attention architectures allow site-specific appearance pathways to adapt to local clothing distributions without perturbing the globally shared identity encoder. The governance implications of federated identity systems are profound, as the absence of a central biometric database reduces the risk of mass data breaches yet complicates centralized audit trails required by accountability frameworks.

The persistence of identity pseudonyms across sessions and sites is another deployment dimension that bridges technical architecture and governance. While transformer-based systems can generate discriminative embedding vectors for each person, the mapping between these embeddings and long-term identity labels must be carefully managed to comply with data minimization principles. Short-lived, encounter-based pseudonyms that are rotated periodically can facilitate attribute analytics without enabling longitudinal tracking, but they require the transformer decoder to be stateless with respect to identity across pseudonym boundaries. Implementing such constraints at the infrastructure level demands close collaboration between system architects and privacy engineers from the earliest design phases, rather than retrofitting legal compliance onto fully specified systems.

6. Governance, Fairness, and Policy Implications

The deployment of multimodal human attribute recognition and identity matching systems at scale raises a constellation of governance challenges that extend well beyond standard performance metrics. Demographic fairness concerns are particularly acute: studies of commercial facial analysis systems [19] have shown that error rates can vary significantly across gender and skin-type intersections, and analogous biases have been documented in person re-identification pipelines [20] where certain clothing styles, body shapes, or gait patterns systematically yield lower matching accuracy. Multimodal fusion can either mitigate or amplify these biases, depending on how modality-specific features are weighted during training. For instance, if the training data over-represents infrared signatures of one demographic group due to sensor placement biases, the cross-attention mechanism may learn to rely predominantly on that modality for that group while using RGB for others, creating a fragmented and inequitable decision surface. Governance frameworks must therefore mandate disaggregated evaluation across protected attributes and require audit trails that log modality contribution weights for contested identity matches.

The global expansion of AI-driven surveillance [21] has prompted regulatory responses that directly constrain the design space of identity inference systems. The European Union's General Data Protection Regulation classifies biometric data used for identification as a special category requiring explicit consent and stringent safeguards. Transformer architectures that fuse multiple biometric modalities may fall under even stricter oversight if the combined data stream enables unique identification where no single modality alone would suffice. Moreover, the ability of multimodal systems to infer sensitive attributes—health status from gait abnormalities, emotional state from facial thermal patterns—expands the scope of potentially intrusive analytics, pushing against the principle of contextual integrity [22] that

holds information flows should be appropriate to the context in which they were collected. Addressing these concerns proactively requires embedding policy-aware gating mechanisms within the transformer decoder that can restrict attribute inference based on the declared purpose of each camera stream and the jurisdiction’s legal framework.

Fairness-aware training of multimodal transformers involves more than balancing demographic subgroups in the training set; it necessitates rethinking the loss functions, attention regularizers, and fusion strategies themselves. Decoupled feature pathways, as described in [11], offer a structural mechanism for fairness because they separate identity-related signals from appearance factors that may be correlated with protected attributes. By regularizing the mutual information between these pathways, the system can be encouraged to base identity decisions on geometric and motion patterns that are less entangled with race and gender than texture and color cues. However, such regularization must be carefully calibrated: overly aggressive decoupling can inadvertently penalize legitimate cross-modal correlations that are critical for accuracy in rare pose conditions, effectively introducing a different form of fairness-accuracy trade-off. Continuous monitoring through independent bias audit services, coupled with versioned model registries that track the impact of each update on subgroup performance, is essential for maintaining accountability over the system’s lifecycle.

Transnational deployment of identity recognition infrastructure further complicates governance because legal standards for biometric collection, retention, and cross-border transfer are fragmented. A system designed to meet GDPR requirements in Europe may be non-compliant with sectoral privacy laws in other regions. Architecting the transformer pipeline such that modality processing and identity linking are encapsulated in geographically bounded microservices, with cryptographically enforced data residency policies, provides a technical substrate for regulatory interoperability. Auditability, transparency reporting, and human-in-the-loop override mechanisms become first-class architectural requirements rather than afterthoughts. Engaging with civil society organizations and affected communities during the requirements phase helps align system capabilities with societal expectations, reducing the risk of public backlash and regulatory sanction that can derail large-scale infrastructure projects.

7. Conclusion

Transformer-based multimodal human attribute recognition and identity matching represent a significant advance in the capacity of intelligent infrastructures to operate reliably under the messy, variable conditions of the real world. By weaving together complementary sensing modalities through carefully structured attention mechanisms, these systems can transcend the brittleness that has long limited person re-identification to controlled environments. However, the technical sophistication of cross-modal transformers must be matched by an equally sophisticated engagement with the systems-level trade-offs that govern their deployment, including edge-cloud resource allocation, model sustainability, privacy-preserving learning, and fairness verification. Decoupled attention architectures, particularly those that separate feature-driven and multimodal fusion pathways, provide a promising blueprint for building models that are simultaneously robust to appearance variations, adaptable to evolving sensor networks, and structurally amenable to fairness audits.

Looking ahead, the research community must move beyond single-task leaderboard optimization and embrace a multidisciplinary paradigm in which architectural innovation is shaped by legal, ethical, and operational constraints from the outset. This entails developing standardized benchmarks that capture not only matching accuracy but also energy

consumption per query, subgroup fairness variance, modality failure resilience, and compliance with data minimization principles. It also requires nurturing tighter collaboration between computer vision researchers, systems engineers, human-computer interaction specialists, and policy scholars to ensure that the next generation of identity inference systems enhances public safety and service personalization without eroding the autonomy, dignity, and privacy of the individuals they observe. Only through such holistic, infrastructure-conscious scholarship can the promise of multimodal transformer-based human analysis be realized in a manner that is technically robust, socially legitimate, and institutionally sustainable.

References

1. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems* (pp. 5998–6008).
2. Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., Uszkoreit, J., & Houlsby, N. (2021). An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*.
3. Ye, M., Shen, J., Lin, G., Xiang, T., Shao, L., & Hoi, S. C. H. (2021). Deep learning for person re-identification: A survey and outlook. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(6), 2872–2893.
4. Liu, X., Zhao, H., Tian, M., Sheng, L., Shao, J., Yan, S., & Wang, X. (2017). Hydraplus-net: Attentive deep features for pedestrian analysis. In *Proceedings of the IEEE International Conference on Computer Vision* (pp. 350–359).
5. Wu, A., Zheng, W. S., Yu, H. X., Gong, S., & Lai, J. (2017). RGB-infrared cross-modality person re-identification. In *Proceedings of the IEEE International Conference on Computer Vision* (pp. 5380–5389).
6. He, S., Luo, H., Wang, P., Wang, F., Li, H., & Jiang, W. (2021). Transreid: Transformer-based object re-identification. In *Proceedings of the IEEE/CVF International Conference on Computer Vision* (pp. 15013–15022).
7. Wu, Z., Huang, Y., Wang, L., Wang, X., & Tan, T. (2016). A comprehensive study on cross-view gait based human identification with deep CNNs. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 39(2), 209–226.
8. Li, S., Bak, S., Carr, P., & Wang, X. (2018). Diversity regularized spatiotemporal attention for video-based person re-identification. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (pp. 369–378).
9. Li, Y., Yao, H., Duan, L., & Gong, Y. (2019). Disentangled representation learning for person re-identification. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 5916–5925).
10. Zhang, Q., Lai, J., Feng, Z., Xie, X., & Huang, Q. (2022). Part-aware transformer for visible-infrared person re-identification. In *Proceedings of the AAAI Conference on Artificial Intelligence* (pp. 3429–3437).

11. Ding, Y., Wang, X., Yuan, H., Qu, M., & Jian, X. (2025). Decoupling feature-driven and multimodal fusion attention for clothing-changing person re-identification. *Artificial Intelligence Review*, 58(8), 241.
12. Liu, J., Zha, Z. J., Chen, D., Hong, R., & Wang, M. (2020). Attribute-driven feature disentanglement and temporal interaction learning for video-based person re-identification. *IEEE Transactions on Image Processing*, 29, 4756–4770.
13. Li, X., Zheng, W. S., Wang, X., Xiang, T., & Gong, S. (2020). Long-term cloth-changing person re-identification. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 6430–6439).
14. Jia, J., Ruan, Q., & An, G. (2020). Disentangled person image generation and its application to person re-identification. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 6670–6679).
15. Zhang, T., Liu, L., & Gong, S. (2021). Person re-identification by person-specific adversarial attacks. *IEEE Transactions on Information Forensics and Security*, 16, 702–715.
16. Shi, W., Cao, J., Zhang, Q., Li, Y., & Xu, L. (2016). Edge computing: Vision and challenges. *IEEE Internet of Things Journal*, 3(5), 637–646.
17. Cheng, Y., Wang, D., Zhou, P., & Zhang, T. (2017). A survey of model compression and acceleration for deep neural networks. *IEEE Signal Processing Magazine*, 35(1), 126–136.
18. Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated learning: Challenges, methods, and future directions. *IEEE Signal Processing Magazine*, 37(3), 50–60.
19. Buolamwini, J., & Gebru, T. (2018). Gender shades: Intersectional accuracy disparities in commercial gender classification. In *Proceedings of the Conference on Fairness, Accountability and Transparency* (pp. 77–91).
20. Yu, R., Li, Z., & Tao, D. (2022). Debiased person re-identification. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(10), 6854–6870.
21. Feldstein, S. (2019). The global expansion of AI surveillance. *Carnegie Endowment for International Peace*.
22. Nissenbaum, H. (2004). Privacy as contextual integrity. *Washington Law Review*, 79(1), 119–158.
23. Schwartz, R., Dodge, J., Smith, N. A., & Etzioni, O. (2020). Green AI. *Communications of the ACM*, 63(12), 54–63.