

Concept Paper

Not peer-reviewed version

Concept and Feasibility of Heliocentric Artificial Planets for Scalable Power Generation and Autonomous Space Infrastructure

Jeongsik Choi *

Posted Date: 11 February 2026

doi: 10.20944/preprints202602.0848.v1

Keywords: heliocentric infrastructure; space-based solar power; wireless power transmission; systems-of-systems; autonomous operations; orbital congestion



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](https://creativecommons.org/licenses/by/4.0/), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Concept Paper

Concept and Feasibility of Heliocentric Artificial Planets for Scalable Power Generation and Autonomous Space Infrastructure

Jeongsik Choi

Independent Researcher (Affiliation intentionally withheld); wjdtlrdl97@gmail.com

Abstract

Earth-centric satellite systems are increasingly constrained by orbital congestion, collision exposure that scales nonlinearly with constellation size, and geometry-driven power intermittency. This paper proposes Heliocentric Artificial Planets (HAPs): modular, actively controlled heliocentric hubs that deliver persistent solar power, autonomous coordination, and data aggregation for distributed satellite networks. We provide quantified scaling laws, explicit numerical evaluations, and a system-of-systems architecture that together demonstrate physical feasibility within known laws of orbital mechanics and electromagnetic transmission. The concept reframes future space systems from spacecraft-centric to infrastructure-centric design and positions heliocentric placement as a structural solution to Earth-orbit scalability limits.

Keywords: heliocentric infrastructure; space-based solar power; wireless power transmission; systems-of-systems; autonomous operations; orbital congestion

1. Introduction

Artificial satellites have progressed from single-purpose spacecraft into dense, networked infrastructures that support navigation, communications, Earth observation, scientific missions, and increasingly, real-time economic activity. Yet the scaling of such systems in Earth-centric orbits is not purely an engineering optimization problem; it is constrained by geometry and coupling in a bounded phase space. When tens of thousands of objects occupy similar altitude bands and inclinations, the system-wide effort required to maintain safety, spectrum compatibility, and operational continuity rises faster than the incremental capability delivered. Debris environment studies show that collision exposure grows nonlinearly and that collision-generated fragments can amplify risk in ways that are difficult to reverse once density passes a critical threshold [16–19]. This is especially relevant for large constellations where operational success depends on a reliable “traffic system” for orbit maintenance, conjunction assessment, and end-of-life disposal. The physical environment thus acts as a scaling governor: the more the orbit is populated, the more resources must be diverted to risk management rather than mission output.

In parallel, global energy demand and the need for low-carbon power motivate renewed attention to space-based solar power (SBSP). The conceptual appeal of SBSP is straightforward: sunlight in space is abundant, predictable, and can be harvested above weather systems and aerosol variability. However, conventional SBSP proposals have typically concentrated on geostationary orbit (GEO) or cislunar regimes, which remain subject to eclipse seasons, station-keeping burdens, and regulatory congestion [1–6]. These constraints limit the achievable duty cycle and complicate the mass and thermal design, particularly when the aim is not kilowatt- or megawatt-class payload power but infrastructure-scale delivery. The key point is that Earth-centric placements couple power generation to the Earth–Sun geometry; therefore, even large collector areas cannot avoid eclipse-driven intermittency without substantial storage mass and reliability penalties.

This paper proposes a complementary and infrastructure-centric architecture: Heliocentric Artificial Planets (HAPs). A HAP is not a gravitationally bound world but a modular heliocentric hub that functions as a persistent energy source, communication and computation node, and autonomy coordinator for distributed satellite networks. The proposition is that moving infrastructure-scale functions away from Earth's crowded orbital shells decouples growth in capability from growth in congestion. This is analogous to how terrestrial infrastructure evolved: when cities grew, power generation, logistics hubs, and backbone networks were re-architected rather than simply densified. Here, the heliocentric environment supplies near-continuous solar exposure and a comparatively low-traffic operational domain, enabling linear scaling in energy output with collector area while keeping collision exposure largely independent of expansion. The remainder of this work quantifies these claims with explicit scaling laws, numerical evaluations, and system-of-systems reasoning grounded in established aerospace engineering methods [12,20,24].

2. Structural Scaling Limits in Earth-Centric Orbits

The limitations of Earth-centric architectures become clearest when expressed as scaling relations. Consider an orbital shell containing N active spacecraft. A first-order model for expected collision exposure P_c is proportional to the number of interacting pairs, leading to $P_c \propto N^2 \sigma v$, where σ is an effective cross-sectional area and v is a characteristic relative velocity [16]. This relationship does not assert that collisions are inevitable, but it implies that the number of potentially hazardous conjunctions grows quadratically as density increases. Debris modeling and long-term environment studies corroborate that small perturbations in traffic or disposal compliance can substantially alter the growth of hazardous populations over decadal timescales [17–19]. Importantly, this quadratic trend is structural: it originates from the combinatorics of pairwise encounters within a bounded volume, not from any particular spacecraft technology.

By contrast, many measures of delivered capability scale roughly linearly with satellite count. For example, increasing N can improve temporal coverage or aggregate throughput, but the marginal gain typically diminishes as coverage saturates or as network bottlenecks appear. Hence a divergence emerges: negative externalities scale $\sim N^2$ while benefits scale $\leq N$. Operationally, this divergence shows up as escalating propellant consumption for avoidance maneuvers, increasing burden on space surveillance and tracking networks, and rising system complexity in scheduling and communications. Even with onboard autonomy, spacecraft remain embedded in the same physical density, so autonomy mainly shifts decision-making location rather than removing the coupling. The Kessler-type cascade mechanism further amplifies this concern: a single breakup event can create fragments that increase collision probability for many other objects, potentially producing a feedback process that is costly to arrest once initiated [16].

Orbital congestion also has non-collision aspects. Spectrum allocation in crowded regimes becomes harder as systems proliferate; regulatory coordination and interference mitigation can impose schedule constraints and reduce effective capacity. Moreover, the presence of large numbers of satellites complicates mission assurance for high-value assets, including crewed vehicles and national critical infrastructure. These costs are not always internalized by individual operators, creating a collective-action dynamic that makes purely incremental solutions fragile. From an engineering standpoint, this suggests that Earth-centric orbits are approaching a regime where scaling is dominated by risk management rather than mission output.

The HAP concept responds to this scaling mismatch by relocating infrastructure-scale functions (e.g., bulk energy generation, centralized processing, backbone relays) to a domain where increasing capability does not require increasing local traffic density. In heliocentric placement, expansion can be achieved by adding collector area or modules without requiring the same orbital shell to host more independent spacecraft. Thus, the scaling of collision exposure can be decoupled from the scaling of

delivered infrastructure capacity. This does not eliminate the need for Earth-orbit satellites, but it reassigns their role to edge sensing/actuation while heavy infrastructure resides in a less congested domain. Subsequent sections quantify how this decoupling changes both energy and operations scaling.

3. Power Intermittency and Orbital Geometry Constraints

Power availability is an equally fundamental scaling constraint. Solar arrays deliver power proportional to incident irradiance and projected area, but Earth-centric orbits impose geometric interruptions that cannot be removed by engineering refinement alone. In LEO, eclipse periods occur every orbit for many inclinations, forcing spacecraft to rely on storage for a substantial fraction of time. In GEO, seasonal eclipse periods persist, and although they occupy a smaller fraction of the year, they create a deterministic requirement for storage and thermal management that scales with delivered power. For SBSP architectures, these eclipse constraints are more severe because the goal is continuous baseload-like delivery rather than intermittent supply. If a GEO SBSP platform aims to deliver gigawatt-class power, the storage required to bridge eclipse intervals becomes extremely large, adding mass, cost, and degradation-driven replacement risk. At the subsystem level, battery and power electronics lifetime can become the dominant driver of maintenance cycles, especially under radiation exposure and thermal cycling [28–30].

This geometric reality has motivated decades of SBSP studies exploring alternative placements and architectures. Early concepts emphasized GEO to maintain near-constant line-of-sight to fixed ground receivers, while later studies considered modularity, phased arrays, and different frequency bands to improve safety and efficiency [1–6]. Yet the Earth–Sun geometry remains a hard boundary: in any Earth-bound orbit, the planet can occlude the Sun and produce shadowing, and station-keeping in certain regimes can impose additional power requirements. Even when eclipse durations are manageable, the thermal cycling associated with entering and leaving shadow induces mechanical stress, drives material fatigue, and increases the probability of failure in large deployable structures. These effects are typically tolerable for conventional satellites but become central for infrastructure-scale megawatt to gigawatt systems intended to operate for decades.

Heliocentric placement addresses these constraints at the root. In an Earth-like heliocentric orbit near 1 AU, a platform is illuminated essentially continuously because there is no Earth shadowing. The duty cycle can exceed 99% depending on exact geometry and orientation, and the remaining interruptions can be designed around with small-scale regulation storage rather than eclipse-bridging storage. This changes the design space: storage mass scales with transient regulation needs rather than with worst-case eclipse energy. It also improves thermal stability because the system no longer experiences repeated full-scale transitions between sunlight and shadow. In turn, improved thermal stability reduces structural fatigue and can increase effective lifetime of photovoltaic and thermal subsystems.

Therefore, power intermittency in Earth-centric orbits is not merely an inconvenience; it is a structural constraint that pushes SBSP systems toward heavier, more complex designs with shorter effective lifetimes. A heliocentric hub offers a structural pathway to continuous energy generation, which then enables the broader HAP role: persistent power for onboard computing and communication, continuous support for satellite edge nodes, and stable operation as a backbone infrastructure. The next sections formalize this advantage within a broader trilemma that also accounts for traffic density and autonomy.

Comparative Advantage over Earth–Moon Lagrange Points (L1/L2)

Earth–Moon Lagrange points, particularly L1 and L2, offer clear advantages in terms of proximity to Earth and reduced communication latency. As a result, they have been widely adopted for scientific observatories and limited deep-space missions. However, when evaluated from the perspective of infrastructure-scale deployment, L1/L2 locations exhibit fundamental gravitational and geometric limitations that severely constrain long-term scalability.

By definition, Lagrange points represent local equilibrium solutions of the restricted three-body problem and are not volumetric regions but effectively point-like locations surrounded by narrow families of quasi-periodic halo or Lissajous orbits. These orbits are inherently unstable, requiring continuous station-keeping to counteract exponential divergence driven by solar and terrestrial perturbations. As the number and physical extent of co-located assets increase, mutual gravitational interactions and control coupling intensify, leading to rapidly escalating propellant consumption and operational complexity.

In contrast, the heliocentric orbit proposed in this study does not rely on a discrete equilibrium point but instead exploits the full circumsolar orbital manifold at approximately 1 AU. This distinction is critical: a heliocentric configuration provides effectively unbounded azimuthal deployment space, allowing large-scale artificial planetary infrastructure to be distributed along an extended orbital arc or ring without introducing strong mutual gravitational coupling. As a result, infrastructure scaling is governed primarily by manufacturing and assembly capacity rather than orbital crowding or station-keeping constraints.

From an environmental standpoint, L1/L2 regions reside within complex interaction zones shaped by the Earth's magnetotail, solar wind variability, and transient geomagnetic disturbances. These effects introduce non-stationary thermal and radiation conditions that complicate the long-term operation of large-area power collection structures. A heliocentric orbit, by contrast, offers a predictable and quasi-steady solar radiation environment with minimal shadowing and reduced thermal cycling. This stability significantly enhances structural survivability, thermal control efficiency, and degradation predictability over multi-decade timescales.

Consequently, while L1/L2 architectures remain well suited for point-mission observatories and Earth-proximal deep-space gateways, they are intrinsically ill-matched to the demands of planetary-scale energy and coordination infrastructure. The heliocentric orbital regime uniquely satisfies the combined requirements of spatial scalability, dynamical stability, and environmental predictability, thereby forming the foundational justification for the Heliocentric Artificial Planet concept proposed in this work.

4. Energy–Traffic–Autonomy Trilemma

The feasibility of scaling space infrastructure can be expressed as a trilemma between continuous energy availability (E), traffic density (T), and operational autonomy (A). In practical terms, E captures whether the architecture can deliver near-continuous power for payloads or for export; T captures the extent to which the operational regime is crowded, raising conjunction and coordination burden; and A captures how much decision-making and fault management can be executed without constant human-in-the-loop operations. Earth-centric systems struggle to maximize all three simultaneously because each variable couples to the others through physical and operational constraints. In LEO, high T is attractive for low latency and short link budgets, but it imposes severe constraints on E through the N^2 scaling discussed earlier [16–19]. Maintaining A in such regimes requires expensive onboard autonomy and high-fidelity ground surveillance, and even then the environment remains coupled. In GEO, T is lower but E is limited by eclipse seasons and by platform mass growth, and the cost of station-keeping and space situational awareness remains nontrivial. In Earth–Sun L1/L2 regimes, E can be high for certain orientations and science missions, but the orbits are metastable and require station-keeping; moreover, these regions are not ideal for exporting large continuous power to Earth due to geometry and distance considerations [13–15].

This trilemma becomes particularly salient when the mission shifts from “a satellite does X” to “space infrastructure continuously provides X.” Infrastructure implies persistent service, predictable performance envelopes, and the capacity to expand without catastrophic growth in risk. Terrestrial analogies are instructive: data centers are not built inside crowded traffic intersections; power plants are not placed in places where logistics are irreducibly congested. In space, however, Earth-centric architectures often try to do exactly that: they place the bulk of infrastructure functions (power generation, processing, coordination) inside the most congested orbital shells because those shells are

closest to Earth. The result is a scaling deadlock where each attempt to increase E (bigger platforms) or increase A (more autonomy) tends to raise T-related complexity (more mass, more cross-section, more coordination).

A heliocentric artificial planet expands the feasible region because it changes the coupling relationships. In heliocentric placement, T can be kept low even as E grows, since capability scaling is achieved primarily by adding area and modules rather than by adding independent objects to a crowded shell. Moreover, because a HAP functions as a hub, A can be implemented as shared infrastructure: high-performance computing and AI control can be centralized and updated, rather than replicated across every edge satellite. The trilemma thus shifts from a “hard triangle” to a “relaxed feasible region”: E can be increased linearly with collector area, T remains low by design, and A becomes more tractable because coordination and fault management can be centralized. This does not eliminate operational challenges—e.g., long-range power/data transmission—but it changes the scaling laws that dominate risk and cost. The remainder of the paper uses explicit equations and numerical examples to show how heliocentric placement enables this architectural rebalancing.

5. Definition and System Architecture of Heliocentric Artificial Planets

A Heliocentric Artificial Planet (HAP) is defined here as a modular, actively controlled heliocentric infrastructure that is “planet-like” only in the sense that it follows a stable heliocentric trajectory, not in mass, gravity, or habitability. The HAP is conceived as an engineered system-of-systems whose fundamental purpose is to provide persistent energy, communication, computation, and coordination services to a distributed set of space assets. This definition is intentionally engineering-centric: the HAP is not a single monolithic spacecraft but a growing infrastructure assembled from repeated modules, each within the mass and volume class of large satellites or space station elements. This modularity enables phased deployment and replacement, which is crucial for avoiding the “single-point-of-failure megastructure” critique that often undermines large space concepts.

Architecturally, the HAP can be decomposed into layered subsystems: (i) a solar collection layer consisting of photovoltaic arrays and/or solar-thermal collectors sized for gigawatt-class energy capture; (ii) a power bus and storage layer that provides regulation, conversion, and distribution; (iii) an AI control and communication core that performs state estimation, fault diagnosis, resource scheduling, and beam/relay management; and (iv) interface subsystems that connect the HAP to satellite edge nodes and to Earth-based receivers through relay networks. Each layer can be expanded independently, enabling incremental growth in capability without requiring a complete redesign. The system-of-systems framing aligns with established engineering principles: complex infrastructures are best built from loosely coupled subsystems with clear interfaces, enabling upgrades, redundancy, and resilience [24,25].

A crucial design goal is service continuity. Because the HAP is intended as infrastructure, it must tolerate partial degradation and continued operation during maintenance. This implies both redundancy and reconfigurability at the module level: power routing must adapt to failed strings; communication and beamforming must tolerate element loss; computing must operate under graceful degradation. These characteristics are similar to terrestrial data center and grid architectures, where fault-tolerant design is achieved through redundancy, modular replacement, and automated monitoring. The HAP is thus a space analog of a hybrid energy plant plus network backbone plus autonomous operations center.

Finally, the HAP architecture explicitly anticipates a hierarchical relationship with conventional satellites. Rather than replacing satellites, the HAP reduces the burden on them by offloading high-power functions (bulk energy generation, backbone processing, long-haul relay management) to the heliocentric hub. Satellites become edge sensors and actuators that can be simpler, cheaper, and more numerous without proportionally increasing infrastructure-scale complexity. This role separation is central to the HAP rationale: it is the mechanism by which scaling laws improve. Subsequent sections

quantify heliocentric placement, energy capture, transmission, and the operational implications of this architecture.

6. Orbital Mechanics and Heliocentric Stability

Heliocentric placement is attractive not because it is exotic, but because it provides a stable, low-traffic environment with near-continuous solar exposure. A practical baseline is an Earth-like heliocentric orbit with a small phase offset (e.g., leading or trailing Earth by a few degrees). Such a configuration preserves an orbital period close to one year, simplifies relative geometry for relay planning, and keeps solar irradiance approximately constant. The orbital mechanics of heliocentric trajectories are well understood: at 1 AU, the Sun's gravity dominates, and perturbations from planets act as secondary effects that can be managed with occasional station-keeping. Compared to Earth–Sun L1/L2 libration point orbits, which are metastable and require continuous control to remain within bounded regions of phase space, a heliocentric orbit is inherently stable in the two-body approximation and requires comparatively lower intervention for long-term maintenance [12–15].

The main operational consideration for an Earth-phase-offset orbit is maintaining the desired relative angle over time. Small differences in semi-major axis or orbital energy can cause drift. However, because the HAP is infrastructure-scale, modest propulsion budgets distributed across many modules can provide periodic corrections. The required Δv depends on the chosen orbit design and perturbation environment, but the key point is qualitative: maintaining a heliocentric orbit near 1 AU does not demand continuous station-keeping comparable to that required by halo or Lissajous orbits around libration points. This reduces the propellant fraction allocated to orbit maintenance and improves the feasibility of multi-decade operation.

From a thermal standpoint, the absence of eclipse events yields a stable illumination profile that simplifies thermal control. The HAP can be designed with constant orientation strategies (e.g., sun-pointing collectors plus articulated radiators) without the need to handle frequent transitions through shadow. Stable thermal conditions reduce fatigue, reduce the amplitude of thermoelastic deformation, and improve long-term alignment stability for phased arrays and large structures—an important consideration for wireless power transmission beamforming.

From a traffic standpoint, heliocentric space near 1 AU is not empty, but it is far less crowded than LEO and GEO. The probability of conjunction with Earth-orbit debris is effectively eliminated, and collision risk is dominated by the system's own internal geometry and by rare heliocentric objects. This environment is thus better suited for scaling large structures whose primary constraint is surface area and alignment rather than proximity operations in a crowded shell. In short, heliocentric placement changes the governing constraints from orbital congestion and eclipse geometry to manufacturing, assembly, and long-range transmission—constraints that are challenging but scale more favorably for infrastructure. The next section leverages these orbital benefits to quantify energy capture and demonstrate explicit power numbers at realistic efficiencies.

7. Solar Power Generation and Quantitative Scaling

The energy advantage of heliocentric placement can be captured with a small set of equations that yield transparent numerical results. At 1 AU, the average solar irradiance (solar constant) is approximately $S_0 \approx 1361 \text{ W}\cdot\text{m}^{-2}$ [1,2]. If the HAP employs photovoltaic conversion and associated power conditioning with a system-level efficiency η , the usable electrical power density is $P_u = \eta S_0$. The system-level efficiency is the relevant parameter because it accounts not only for cell conversion but also for packing factor, wiring losses, maximum power point tracking, conversion overhead, and degradation margin. This value can be interpreted as “continuous watts per square meter of collector” when illumination is continuous and thermal conditions are well regulated.

Total continuous power is then $P = P_u A$, where A is the effective collector area. Substituting $A = 10 \text{ km}^2 = 1 \times 10^7 \text{ m}^2$ yields $P \approx 408 \times 10^7 \approx 4.08 \times 10^9 \text{ W}$, i.e., about 4.08 GW. The linearity here is crucial: doubling area doubles power, with no intrinsic geometric penalty. For comparison, if an Earth-centric

platform experiences a duty cycle D due to eclipses (for example, $D = 0.85$), the average delivered power becomes $P_{av} = D \cdot P_u A$, meaning that achieving the same average power requires A increased by $1/D$, along with storage sized for $(1-D)$ intervals. In heliocentric operation with $D \approx 1$, storage requirements are reduced to transient regulation, not eclipse bridging.

Mass scaling depends on structural areal density μ ($\text{kg}\cdot\text{m}^{-2}$), which includes panels, truss, deployment, wiring, and thermal. If μ is $2 \text{ kg}\cdot\text{m}^{-2}$ —a demanding but not implausible long-term target for large deployables—then a 10 km^2 collector implies a structural mass of $2 \times 10^7 \text{ kg} = 20,000$ metric tons. This is large by today's standards but can be approached through modular assembly and long-horizon industrialization. Crucially, the mass-to-power ratio becomes $(\mu/P_u) \approx (2)/(408) \approx 0.0049 \text{ kg}\cdot\text{W}^{-1}$, or about $4.9 \text{ kg}\cdot\text{kW}^{-1}$, which is within the conceptual envelope of SBSP feasibility studies that assume significant reductions through mass production and in-space assembly [2–6]. If μ is $4 \text{ kg}\cdot\text{m}^{-2}$, the ratio doubles, illustrating the importance of structural and manufacturing innovations. These calculations do not prove near-term implementation, but they do demonstrate that no physical law prohibits gigawatt-class continuous generation with plausible efficiency and mass density assumptions. The next challenge is delivering that power and associated data services to useful endpoints, which is addressed via wireless transmission architectures.

8. Wireless Power Transmission Architecture and Numerical Link Budgets

Power export from a heliocentric hub is often criticized on the basis of distance, but the relevant conclusion depends on architecture. A naïve single-hop transmission from 1 AU to Earth is indeed challenged by free-space path loss. The first-order received power for a line-of-sight link can be expressed with the Friis relation: $P_r = P_t G_t G_r (\lambda/4\pi R)^2$, where P_t is transmitted power, G_t and G_r are antenna gains, λ is wavelength, and R is separation distance [7–9]. Consider a microwave frequency of 2.45 GHz, for which $\lambda \approx 0.122 \text{ m}$. Suppose the HAP transmits $P_t = 5 \text{ GW}$ (a fraction of the 10 km^2 generation example). If both the transmitter and receiver employ phased arrays with gains $G_t = G_r = 10^7$ (~70 dBi), which corresponds to very large effective apertures, then at $R = 1 \text{ AU} \approx 1.496 \times 10^{11} \text{ m}$, the factor $(\lambda/4\pi R)^2$ becomes extraordinarily small, yielding P_r far below useful levels. This numerical result is correct—and it is precisely why realistic architectures do not attempt single-hop delivery.

The structural solution is distance segmentation via relay nodes. If the effective transmission path is divided such that each hop occurs over distances on the order of 10^7 – 10^8 m (e.g., GEO scale $\sim 4 \times 10^7 \text{ m}$, or cislunar distances), then the same gains yield P_r improved by $(R_{\text{old}}/R_{\text{new}})^2$. Reducing distance by a factor of 10^4 increases P_r by a factor of 10^8 . In practice, relay nodes can refocus beams, perform frequency conversion, and manage safety interlocks. This concept aligns with decades of SBSP studies that emphasize phased arrays, rectennas, and controlled power density as enabling features [8–11,38]. The engineering conclusion is that heliocentric distance does not forbid power export; it demands a hierarchical architecture where the heliocentric hub provides generation and first-hop beaming to strategically placed relays, which then deliver to Earth or to Earth-orbit receivers.

Safety and regulatory constraints further motivate relays. A beam intended for Earth must satisfy power density limits to avoid hazards to aviation, wildlife, and human exposure. Using relays permits smaller spot sizes where appropriate, dynamic beam steering with exclusion zones, and the ability to terminate transmission rapidly if pointing deviates. Phased array control technologies for beam shaping and sidelobe management are a known research area, and the HAP concept leverages rather than invents this foundation [9,10]. Laser transmission is an alternative with higher power density but stricter atmospheric and eye-safety concerns, and it may be most attractive for space-to-space transfer or as a complementary channel for data and precision power routing [37].

Finally, the same relay architecture supports data aggregation and communication. High-gain links used for power can also host high-capacity communications, enabling the HAP to function as a backbone node for satellite telemetry and for deep-space relay. This dual-use of aperture and pointing control improves overall system economics. The subsequent sections address the environmental

reliability of continuous heliocentric operation and the autonomy framework required to manage such a complex, distributed infrastructure.

9. Thermal, Radiation, and Degradation Environment

Long-term feasibility requires that a HAP survive the heliocentric environment for decades with manageable degradation. The thermal environment in heliocentric orbit near 1 AU is, in some ways, simpler than in Earth-centric orbits because illumination is continuous and predictable. The absence of repeated eclipse cycles reduces large-amplitude temperature swings, which are a major driver of thermo-mechanical fatigue in deployable structures. Reduced cycling can improve the longevity of adhesives, joints, and composite materials, and it can stabilize the alignment of large apertures used for beamforming. However, continuous illumination also implies continuous heating; therefore radiator sizing, emissivity stability, and thermal control architecture remain central. The benefit is that the system can be designed around a steady-state balance rather than around extreme transient cases. This can reduce peak thermal stress and simplify control laws.

Radiation exposure in heliocentric space includes galactic cosmic rays, solar energetic particles, and trapped radiation is largely absent compared to Earth's belts (depending on exact orbit). Electronics must be designed for total ionizing dose, displacement damage, and single-event effects. The relevant mitigation toolbox is well established: shielding, part selection, error-correcting codes, redundancy, and fault-tolerant computing. Radiation effects in microelectronics have been extensively characterized, and design methodologies exist for ensuring long-duration performance even under harsh conditions [28,29]. Models for cumulative heavy-ion deposition and dose accumulation support mission-level predictions and margins [30].

A key additional factor is space weather: coronal mass ejections and solar energetic particle events can produce transient increases in particle flux that stress electronics and, in extreme cases, solar arrays. Yet this risk is not unique to HAPs; interplanetary spacecraft routinely operate under similar conditions. Moreover, the HAP's role as a heliophysics monitor can provide early warning to Earth-orbit systems, creating a net benefit for the broader space ecosystem.

Degradation of photovoltaic performance over decades is expected due to radiation and micrometeoroid impacts. Modular design addresses this by enabling replacement of degraded arrays without decommissioning the entire platform. The same modularity supports incremental upgrades: as solar cell efficiencies improve, new modules can be added to increase power density without altering the core architecture. This is analogous to how terrestrial power grids integrate new generation technologies while maintaining legacy infrastructure. Thus, the long-term environment does not invalidate the concept; it emphasizes the need for modular replacement and for autonomy that can manage gradual degradation. The next section focuses on autonomy and AI integration as essential enablers for operating such a distributed system without prohibitive human operations cost.

10. Autonomous Operations and AI Integration

A HAP is, by design, an infrastructure-scale system: it will contain thousands of modules, multiple relay links, phased-array elements, and continuously operating power electronics. Operating such a system with traditional ground-centric command and control would be cost-prohibitive. Therefore, autonomy is not an optional enhancement but a core requirement. The relevant autonomy functions include state estimation across distributed subsystems, fault detection and isolation, predictive maintenance scheduling, and optimization of power and communication routing. AI methods are attractive here not as an abstract trend but because they offer scalable approaches for anomaly detection, pattern recognition in telemetry, and decision support under uncertainty. Recent literature highlights the growing role of autonomy and AI in space systems for resilience and operational efficiency [22,23].

A practical autonomy architecture for a HAP is hierarchical. At the module level, local controllers manage safe operation of power converters, thermal loops, and mechanical actuators. At

the cluster level, controllers coordinate groups of modules to provide stable bus voltage, manage energy storage, and maintain structural attitude. At the system level, an AI-enabled supervisory layer performs planning and optimization across the entire platform: it predicts degradation trends, schedules maintenance windows, and allocates beamforming resources to satisfy competing objectives (e.g., exporting power, supporting satellite communications, maintaining safety constraints). This hierarchical scheme mirrors successful terrestrial infrastructure control architectures, where local control ensures safety and stability while supervisory control optimizes performance.

Autonomy also reduces latency and improves safety for power beaming. Beamforming arrays require continuous pointing and sidelobe management; when the platform is exporting large power, deviation events must be detected and corrected rapidly. These requirements are fundamentally faster than human-in-the-loop response times and thus motivate automated control. Similarly, relays can autonomously manage link acquisition, handovers, and redundancy, enabling continuous service even under partial failures.

Because the HAP is modular, autonomy enables “self-healing” at the infrastructure level: failed modules can be isolated electrically, bypassed, and scheduled for replacement. Predictive maintenance can be driven by telemetry trends, such as increasing converter ripple, rising thermal resistance, or changes in array I–V characteristics. These methods convert catastrophic failures into managed degradations, improving availability and reducing risk. The operational impact is substantial: the HAP can function as a stable backbone even while individual components degrade, analogous to how data centers tolerate server failures without service collapse.

Finally, autonomy extends beyond the HAP itself to its relationship with satellite edge nodes. The HAP can serve as an autonomous operations hub that aggregates satellite telemetry, distributes time and ephemeris products, and manages coordinated commanding strategies. This bridges into the next section: a HAP-centric hierarchy that reorganizes conventional satellites into an edge layer supported by a heliocentric infrastructure hub, improving global scalability and resilience.

11. Satellite–HAP Hierarchical Architecture

A central claim of this paper is that HAPs do not replace satellites; they change what satellites are best used for. In a HAP-centric architecture, satellites become edge nodes—distributed sensors and actuators optimized for proximity to targets (Earth, atmosphere, oceans, or deep-space regions of interest). The infrastructure-heavy functions—bulk power generation, high-performance compute for coordination, backbone communications, and long-duration autonomy services—are shifted to the heliocentric hub. This role separation is analogous to cloud–edge computing in terrestrial networks, where edge devices collect data and execute local actions while centralized infrastructure performs heavy computation and coordination. Systems-of-systems engineering recognizes such decompositions as essential for scalable architectures because they reduce tight coupling and permit independent evolution of subsystems [24,25].

Practically, this means that future constellations could be designed with simpler satellites that rely on hub services when available. For example, satellites could offload computationally intensive tasks (e.g., multi-sensor fusion, global scheduling, anomaly triage) to the HAP. Commanding and telemetry aggregation could be centralized, reducing the number of Earth-based ground station passes required for routine operations. The HAP could provide a stable timing reference, ephemeris updates, and routing services for inter-satellite networks. In deep-space contexts, the HAP could function as an intermediate relay and coordination node, reducing dependence on Earth-based networks and improving responsiveness for autonomous missions.

The hierarchy also improves resilience. In today’s architectures, ground networks are a bottleneck: they must service growing constellations, and outages or congestion can propagate into mission failures. A heliocentric hub provides an alternate coordination layer that is less vulnerable to localized terrestrial disruptions. Moreover, because the HAP can be designed with substantial redundancy and continuous power, it can maintain services during events that disrupt Earth-centric

operations (e.g., regional outages, extreme weather affecting ground stations). This perspective frames HAPs as “space ground stations” or “space infrastructure nodes” rather than as spacecraft in the traditional sense.

An important question is whether long-range links impose unacceptable latency or bandwidth constraints. While heliocentric distances increase propagation delay, many coordination tasks are not latency-critical at the millisecond level; they operate on minutes-to-hours timescales. For tasks that require low latency, edge satellites continue to operate locally. The HAP mainly provides strategic and backbone functions that benefit from stability and scale. This division of labor is key to making the architecture credible: it does not demand that every function be centralized, only that infrastructure-scale functions that suffer in Earth-centric scaling be relocated to a better domain.

In summary, the HAP-centric hierarchy is a systems engineering response to scaling limits. It enables satellites to proliferate as edge nodes without requiring infrastructure-scale growth within congested Earth-orbit shells. The next sections address economic scaling, deployment pathways, and the policy and safety frameworks required to realize such an architecture responsibly.

12. Economic Scaling and Deployment Strategy

The economic feasibility of HAPs is primarily a question of scaling and learning curves rather than a question of physical possibility. Infrastructure-scale systems in space are expensive under today’s launch costs and manufacturing paradigms. However, the history of terrestrial infrastructure indicates that large systems become feasible when demand, industrial capacity, and standardization converge. HAPs align with this pattern because they are modular and can be deployed incrementally. Instead of requiring a single massive investment to build a complete gigawatt-class platform, an initial HAP could be launched as a smaller demonstrator providing tens to hundreds of megawatts to space-based consumers (e.g., cislunar missions, propellant depots, high-performance compute nodes). As reliability and operational procedures mature, collector area and transmission capacity can be expanded.

A useful framing is levelized cost of energy (LCOE) over a multi-decade lifetime. While precise LCOE requires detailed cost models, several qualitative advantages are clear: (i) the fuel cost is effectively zero after deployment; (ii) continuous operation increases capacity factor relative to terrestrial solar; (iii) modular replacement extends system life; and (iv) the infrastructure can serve multiple markets (energy export, communication relay, autonomy services), improving revenue diversity. OECD analyses of the space economy highlight the growing strategic and commercial role of space infrastructure and the importance of policy stability for long-term investments [31]. These considerations suggest that early HAPs may be justified not purely on energy cost grounds but also on strategic resilience and service multiplexing.

Deployment strategy should prioritize risk reduction. Phase 1 can focus on Earth-orbit SBSP demonstrations and on autonomous assembly techniques. Phase 2 can develop relay architectures and beam safety interlocks in cislunar space. Phase 3 can deploy small heliocentric hubs that serve as deep-space relays and energy providers for interplanetary missions. Phase 4 can scale to multi-gigawatt platforms with Earth-directed delivery. This phased approach matches the modular nature of HAPs and avoids the “all-or-nothing” risk profile that has historically hindered large SBSP proposals.

Manufacturing is likely to be a dominant cost driver. Therefore, long-term feasibility is tied to industrial methods such as in-space assembly, robotic servicing, and potentially in-space manufacturing of structural elements. Even without lunar materials, the ability to assemble large structures from standardized modules can drive cost down through mass production. The satellite industry already demonstrates how high-volume production changes economics; applying similar principles to HAP modules is plausible if demand and policy align. Autonomy further reduces operational cost, making large systems economically tractable by reducing the workforce required for continuous supervision.

In summary, economic feasibility is not immediate, but neither is it speculative. The HAP architecture is designed to take advantage of learning curves: each deployed module provides incremental capability and operational learning, enabling a pathway where scale emerges gradually. The next section expands on applications and impact, emphasizing why such an infrastructure would be valuable beyond energy export alone.

13. Applications and Societal Impact

HAPs are most compelling when viewed as multi-function infrastructure rather than as single-purpose power stations. Space-based solar power is a central driver, but the same platform that generates gigawatt-class power can also provide backbone communication and computation services. One major application is persistent heliophysics monitoring. Because the HAP is located in heliocentric space and can host large apertures and sensors, it can observe solar activity and detect space weather events that threaten both terrestrial grids and satellite systems. Early warning and coordinated response can reduce disruptions, and the HAP can redistribute its own resources to protect critical subsystems during extreme events. This dual role—energy infrastructure plus space weather monitor—creates value synergies.

Another application is deep-space communication relay. Current deep-space networks rely heavily on Earth-based assets with limited availability and significant scheduling constraints. A heliocentric hub can provide intermediate relays, improve coverage, and reduce dependence on Earth-based stations for routine operations. This becomes more important as exploration extends to lunar, Martian, and asteroid missions. The HAP can serve as a stable node for navigation and time transfer as well, complementing Earth-based systems. For distributed satellite fleets, the HAP can function as a “space ground station,” aggregating telemetry and distributing command products, thereby reducing the load on terrestrial networks.

From an Earth-orbit sustainability perspective, HAPs offer indirect benefits. By offloading infrastructure-scale functions (power generation and backbone coordination) from Earth-centric shells, HAPs reduce pressure to build ever larger, more complex platforms in congested regimes. This can mitigate congestion growth and complement debris mitigation strategies. Policy discussions emphasize the need for new solutions to orbital debris and for governance frameworks that prevent tragedy-of-the-commons outcomes [32]. A HAP-centric architecture provides a technical pathway aligned with such policy goals: it changes the incentive structure by creating valuable services outside the most congested orbits.

Societal impact also includes energy security. Even if SBSP from heliocentric hubs begins by serving space consumers, it can evolve into Earth-directed delivery for critical infrastructure or remote regions. A resilient, continuous power source that is not affected by terrestrial weather or geopolitics could play a role analogous to strategic reserves. Of course, this introduces governance and safety considerations: power beaming must be regulated and controlled, and international norms must manage the dual-use potential. However, large infrastructures often require governance evolution; the question is whether the technical pathway exists. This paper argues that it does.

Finally, HAPs create a technology pull for innovations that are valuable broadly: high-efficiency photovoltaics, lightweight deployable structures, phased-array beamforming, autonomous operations, and robotic servicing. These technologies spill over into conventional satellite systems, deep-space probes, and even terrestrial industries. Thus, even partial realization of the HAP pathway can generate near-term value. The next section addresses governance, safety, and policy constraints, emphasizing how to frame HAP deployment responsibly within the aerospace community and international institutions.

14. Governance, Safety, and Policy Considerations

Deploying heliocentric infrastructure that exports power and coordinates satellites raises governance and safety issues that must be addressed explicitly to maintain credibility and to enable

adoption. Wireless power transmission to Earth, particularly at large power levels, requires stringent safety mechanisms. Beam pointing, sidelobe control, and termination interlocks must be designed such that unintended exposure remains below established thresholds. Decades of microwave power transmission research provide engineering foundations for beam control and rectenna design, but scaling to gigawatt-class requires policy frameworks for siting, airspace coordination, and liability [7–11]. A relay-based architecture helps: it allows power density to be controlled at intermediate nodes and enables additional safety layers before energy reaches the ground.

Spectrum regulation is another constraint. Using standard ISM bands may be impractical at large scale due to interference and coordination. Dedicated bands and international agreements would be necessary. Moreover, a HAP acting as a coordination hub for satellites implies data governance concerns: who has access, under what rules, and how to prevent misuse. These issues are not unique to HAPs; they mirror terrestrial concerns around critical infrastructure and network backbone providers. Yet the international nature of space makes governance more complex. Works on space policy emphasize that technological solutions must be accompanied by institutional evolution to manage shared environments and strategic competition [40].

A further governance issue is orbital stewardship. Even though HAPs are heliocentric, they interact with Earth-orbit systems via relays and via service provision. Policies should ensure that HAP deployment reduces, rather than increases, congestion pressures. For example, if HAPs enable energy and coordination services, they could support more responsible end-of-life disposal by powering deorbit maneuvers or by providing coordination. Conversely, poorly governed deployment could create new congestion points at relay orbits. Therefore, system design and governance must be co-optimized.

Dual-use considerations are unavoidable. High-power beaming technologies and autonomous coordination infrastructures have potential military relevance. A credible civil and commercial pathway requires transparency, safety standards, and possibly international monitoring regimes. The history of other dual-use infrastructures—nuclear energy, GPS, and satellite imaging—suggests that governance frameworks can enable beneficial use while managing risks, but only if they are integrated early. The aerospace community can contribute by defining technical standards for beam safety, fault tolerance, and interoperability, thereby reducing uncertainty and enabling trust.

Finally, for a journal such as *Aerospace*, it is important to emphasize that governance considerations do not undermine technical feasibility; they define deployment constraints. The purpose of this section is to show that the HAP concept is not a naive engineering dream but a system-level proposal that recognizes social and institutional factors. With this framing, the final section discusses remaining technical uncertainties and synthesizes the feasibility argument: HAPs are physically plausible, structurally motivated by scaling limits, and achievable through phased deployment paired with appropriate governance.

15. Discussion and Conclusions

The central argument of this paper is that heliocentric artificial planets are a structurally motivated response to the quantified limits of Earth-centric satellite systems. Earth-orbit architectures face two hard constraints that worsen with scale: collision exposure grows approximately as N^2 in bounded shells, and power availability is limited by eclipse geometry and thermal cycling. These constraints are not eliminated by incremental improvements in spacecraft technology; they stem from the environment. As a result, future space infrastructure should redistribute functions across orbital regimes to restore favorable scaling behavior. The HAP concept accomplishes this by relocating infrastructure-scale power generation, backbone communication, and autonomy services to heliocentric space near 1 AU, where solar exposure is continuous and local traffic density can remain low.

We provided explicit numerical evaluations to demonstrate feasibility. Using the solar constant $S_0 \approx 1361 \text{ W}\cdot\text{m}^{-2}$ and a system efficiency $\eta = 0.30$, a usable power density $P_u \approx 408 \text{ W}\cdot\text{m}^{-2}$ follows. A collector area of 10 km^2 yields $\approx 4.08 \text{ GW}$ continuous electrical power, scaling linearly with area. Mass

scaling depends on structural areal density μ ; while the required masses are large by present launch standards, modular assembly and long-horizon industrialization align with historical pathways for large infrastructure. The power export challenge is architectural: the Friis relation shows that single-hop delivery from 1 AU is impractical, but relay-based segmentation transforms the link budget by orders of magnitude and is consistent with SBSP research on phased arrays, rectennas, and beam control. These results support the claim that the HAP concept does not violate physical laws; its feasibility rests on engineering scale-up and governance.

The HAP concept also reframes satellite roles. Satellites become edge nodes optimized for sensing and actuation near their targets, while the HAP provides stable hub services: energy, computation, coordination, and long-haul relay management. This cloud–edge analogy is not rhetorical; it is a systems-of-systems design principle for reducing coupling and improving scalability. Autonomy becomes an enabling factor: AI and hierarchical control reduce operational cost and permit the infrastructure to operate continuously under partial degradation. Environmental factors—radiation, micrometeoroids, and space weather—are manageable through established mitigation methods and through modular replacement. Continuous thermal conditions may even improve lifetime relative to eclipse-cycling Earth-orbit platforms.

A critical and intentionally conservative interpretation is that HAPs should be pursued as an incremental pathway. Early phases can serve space customers (cislunar and deep-space users) before Earth-directed delivery is attempted at large scale. Such a pathway supports technology maturation and governance development. The aerospace community can contribute by establishing standards for beam safety, relay interoperability, and autonomous fault management. In conclusion, heliocentric artificial planets represent a credible evolution of aerospace infrastructure. They address structural scaling limits in Earth orbit, enable continuous power and hub-level autonomy, and provide a phased path toward a future where space systems are engineered as resilient infrastructure rather than as collections of independent spacecraft.

References

1. Glaser, P.E. Power from the Sun: Its Future. *Science* 162, 857–861 (1968).
2. Mankins, J.C. *Space Solar Power: The First International Assessment of Space-Based Solar Power*. Wiley (2014).
3. Mankins, J.C. A fresh look at space solar power: New architectures, concepts and technologies. *Acta Astronautica* 41, 347–359 (1997).
4. Landis, G.A. Solar power satellites. *Acta Astronautica* 55, 985–990 (2004).
5. Seboldt, W. et al. European roadmap on space-based solar power. *Acta Astronautica* 61, 123–130 (2007).
6. Sasaki, S. et al. A new concept of solar power satellite: Tethered-SPS. *Acta Astronautica* 60, 153–165 (2007).
7. Brown, W.C. The history of power transmission by radio waves. *IEEE Trans. Microwave Theory Tech.* 32, 1230–1242 (1984).
8. Shinohara, N. *Wireless Power Transfer via Radiowaves*. Wiley (2014).
9. Shinohara, N. Beam control technologies for microwave power transmission. *IEEE Microwave Magazine* 17, 59–73 (2016).
10. Matsumoto, H. Research on solar power satellites and microwave power transmission in Japan. *IEEE Microwave Magazine* 3, 36–45 (2002).
11. Dickinson, R.M. Evaluation of a microwave high-power reception-conversion array. NASA CR-134886 (1976).
12. Vallado, D.A. *Fundamentals of Astrodynamics and Applications*. Microcosm Press (2013).
13. Gómez, G. et al. *Dynamics and Mission Design Near Libration Points*. World Scientific (2001).
14. Farquhar, R.W. The utilization of halo orbits in advanced lunar operations. NASA TR R-346 (1970).
15. Howell, K.C.; Barden, B.T.; Lo, M.W. Application of dynamical systems theory to trajectory design. *J. Astronaut. Sci.* 45, 161–178 (1997).
16. Kessler, D.J.; Cour-Palais, B.G. Collision frequency of artificial satellites: The creation of a debris belt. *J. Geophys. Res.* 83, 2637–2646 (1978).

17. Liou, J.-C.; Johnson, N.L. Risks in Space from Orbiting Debris. *Science* 311, 340–341 (2006).
18. Lewis, H.G. et al. The sensitivity of the space debris environment to small perturbations. *Acta Astronautica* 61, 685–692 (2007).
19. Krisko, P.H. Proper implementation of the 1998 NASA breakup model. *Adv. Space Res.* 25, 457–460 (2000).
20. Wertz, J.R.; Everett, D.F.; Puschell, J.J. *Space Mission Engineering: The New SMAD*. Microcosm Press (2011).
21. Sweeting, M.N. Modern small satellites—Changing the economics of space. *Proc. IEEE* 106, 343–361 (2018).
22. Furfaro, R. et al. Artificial intelligence and autonomy for space systems. *Acta Astronautica* 170, 738–750 (2020).
23. Izzo, D.; Simões, L.F. Autonomous guidance and control in deep space: A review. *Adv. Space Res.* 63, 264–276 (2019).
24. Maier, M.W. Architecting principles for systems-of-systems. *Systems Engineering* 1, 267–284 (1998).
25. Jamshidi, M. *Systems of Systems Engineering*. Wiley (2009).
26. Hughes, T.P. *Networks of Power: Electrification in Western Society, 1880–1930*. Johns Hopkins Univ. Press (1983).
27. Smil, V. *Energy Transitions: Global and National Perspectives*. Praeger (2016).
28. Johnston, A.H. Radiation effects in advanced microelectronics technologies. *IEEE Trans. Nucl. Sci.* 45, 1339–1354 (1998).
29. Schwank, J.R. et al. Radiation effects in MOS oxides. *IEEE Trans. Nucl. Sci.* 55, 1833–1853 (2008).
30. Xapsos, M.A. et al. Model for cumulative solar heavy ion energy deposition. *IEEE Trans. Nucl. Sci.* 47, 2218–2223 (2000).
31. OECD. *The Space Economy at a Glance*. OECD Publishing (2019).
32. Pelton, J.N. New solutions for the space debris problem. *Space Policy* 30, 1–6 (2014).
33. ESA. *Space-Based Solar Power: Collecting Solar Energy in Space for Use on Earth*. ESA Report (2022).
34. Dyson, F.J. Search for Artificial Stellar Sources of Infrared Radiation. *Science* 131, 1667–1668 (1960).
35. Forward, R.L. Starwisp and advanced sail concepts. *J. Br. Interplanet. Soc.* 38, 501–506 (1985).
36. Ćirković, M.M. The thermodynamics of large artificial structures in space. *Astrobiology* 6, 567–572 (2006).
37. Schubert, F.H. Energy transmission by laser beams. *Applied Optics* 12, 1748–1752 (1973).
38. National Space Society. *Space Solar Power Position Paper*. NSS (2017).
39. Tommei, G.; Milani, A.; Rossi, A. Orbit determination of space debris: Admissible regions. *Celest. Mech. Dyn. Astron.* 97, 289–304 (2007).
40. Johnson-Freese, J. Space as a strategic asset. *Astropolitics* 5, 1–18 (2007).

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.