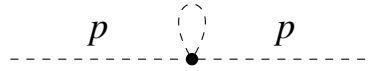


The UV divergences of loop integrals manifest themselves precisely as these poles: in $d = 4 - \varepsilon$ the argument $n - d/2 = n - 2 + \varepsilon/2$ hits a non-positive integer exactly when $n \leq 2$, producing a $1/\varepsilon$ singularity. For example, the tadpole ($n = 1$) needs $\Gamma(1 - d/2) = \Gamma(-1 + \varepsilon/2)$, which has a pole at $\varepsilon = 0$ —the tadpole has superficial degree of divergence 2, and in dimensional regularization this manifests as a $1/\varepsilon$ pole multiplied by m^2 rather than as an explicit Λ^2 . The bubble ($n = 2$) needs $\Gamma(2 - d/2) = \Gamma(\varepsilon/2)$, with a pole encoding the logarithmic divergence. For $n \geq 3$ the argument is strictly positive at $\varepsilon = 0$ and the integral is UV finite.

Since the manipulations involved (Feynman parametrization, momentum shift, reduction to J_n , expansion around $\varepsilon = 0$) are rather mechanical—but instructive on first encounter—we present them here only in outline, deferring the **fully detailed, step-by-step calculations** to Appendix G and the explicit QED computations in Chapter 7.1.

The tadpole: mass renormalization

The one-loop correction to the propagator comes from the tadpole diagram:



Reading the diagram with the Feynman rules: one vertex ($-i\lambda$), one internal propagator ($i/(\ell^2 - m^2 + i\epsilon)$), and a symmetry factor $S = 2$. The origin of this factor is as follows: at the vertex, four ϕ -legs meet. Two of them are connected to the external lines, and the remaining two are contracted with each other to form the loop. Swapping these two internal legs yields exactly the same diagram—the loop propagator starts and ends at the same vertex regardless of which leg we call “first.” This discrete exchange symmetry means the naïve Feynman rules overcount by a factor of 2, so we divide by $S = 2$. Equivalently, the two orientations of the loop momentum (clockwise or counterclockwise) correspond to the same single contraction.¹ This gives:

$$-i\Sigma_{\text{tadpole}} = \frac{-i\lambda}{2} \int \frac{d^4\ell}{(2\pi)^4} \frac{i}{\ell^2 - m^2 + i\epsilon}. \quad (6.6)$$

Power counting. For large $|\ell|$, the integrand behaves as $1/\ell^2$. Since $d^4\ell \sim \ell^3 d\ell$:

$$\int^\Lambda \frac{d^4\ell}{\ell^2} \sim \int^\Lambda \frac{\ell^3 d\ell}{\ell^2} \sim \Lambda^2. \quad (6.7)$$

¹A quick combinatorial check: the vertex carries $(-i\lambda)/4!$, and there are $\binom{4}{2} = 6$ ways to choose which two legs form the loop, times $2! = 2$ ways to assign the remaining two legs to the two external lines, giving 12 distinct contractions. Since $4!/12 = 2 = S$, this confirms $S = 2$.

The integral is **quadratically divergent**.

Evaluation in dimensional regularization. In $d = 4 - \varepsilon$ dimensions, the tadpole integral is a special case of the master integral (6.3) with $n = 1$ and $\Delta = m^2$. Expanding around $\varepsilon = 0$:²

$$\mu^\varepsilon \int \frac{d^d \ell}{(2\pi)^d} \frac{1}{\ell^2 - m^2} = \frac{-im^2}{16\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) - \ln \frac{m^2}{\mu^2} + 1 \right] + O(\varepsilon), \quad (6.8)$$

where $\gamma_E \approx 0.5772$ is the Euler–Mascheroni constant. The $2/\varepsilon$ pole encodes the quadratic divergence that was $\sim \Lambda^2$ in the cutoff language.

The appearance of the logarithm should not be viewed as an artificial trick of the expansion. In dimensional regularization the loop integral is analytically continued away from $d = 4$, so dimensional analysis naturally produces factors such as μ^ε and $(m^2)^{-\varepsilon/2}$. Near $d = 4$ these combine into

$$\left(\frac{\mu^2}{m^2} \right)^{\varepsilon/2} = \exp \left[\frac{\varepsilon}{2} \ln \left(\frac{\mu^2}{m^2} \right) \right] = 1 + \frac{\varepsilon}{2} \ln \left(\frac{\mu^2}{m^2} \right) + O(\varepsilon^2),$$

so the logarithm is the derivative of a power with respect to its exponent, $\lim_{\alpha \rightarrow 0} \frac{x^\alpha - 1}{\alpha} = \ln x$, rather than an arbitrary rewriting. What is scheme-dependent is only which finite combination is subtracted together with the pole—for example, in $\overline{\text{MS}}$ one absorbs the ubiquitous combination $2/\varepsilon - \gamma_E + \ln(4\pi)$ into the definition of the renormalized parameter. Combining now with the vertex factor and the symmetry factor:

$$-i\Sigma_{\text{tadpole}} = \frac{-i\lambda m^2}{32\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) - \ln \frac{m^2}{\mu^2} + 1 \right]. \quad (6.9)$$

The self-energy Σ computed above is a *one-particle irreducible* (1PI) insertion: it cannot be split into two disconnected pieces by cutting a single internal line. The full (dressed) propagator is obtained by **resumming** all possible chains of 1PI insertions on the bare propagator. Denoting the bare propagator by $G_0(p) = i/(p^2 - m^2)$ and the 1PI self-energy by $-i\Sigma$, the dressed propagator is the geometric series

$$G(p) = G_0 + G_0 (-i\Sigma) G_0 + G_0 (-i\Sigma) G_0 (-i\Sigma) G_0 + \cdots = \frac{G_0}{1 + i\Sigma G_0} = \frac{i}{p^2 - m^2 - \Sigma}, \quad (6.10)$$

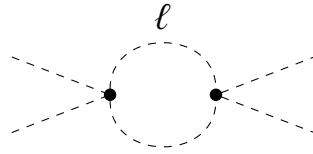
Diagrammatically, this resummation reads:

²The renormalization scale does not disappear: it combines with the factor $(m^2)^{1-\varepsilon/2}$ from (6.3). Indeed, $\mu^\varepsilon (m^2)^{1-\varepsilon/2} = m^2 (\mu^2/m^2)^{\varepsilon/2} = m^2 \left[1 + \frac{\varepsilon}{2} \ln(\mu^2/m^2) + O(\varepsilon^2) \right]$, which is why the final answer contains the term $-\ln(m^2/\mu^2)$.

The key point is that Σ_{tadpole} is UV-divergent, so the *bare* mass m^2 appearing in the Lagrangian (6.1) is *not* the physically measured mass. Instead, the bare mass must itself be divergent, chosen so that the combination $m^2 + \Sigma_{\text{tadpole}}$ yields a finite, measurable quantity m_{phys}^2 . This is the essence of **mass renormalization**: the divergence is not in any observable—it is absorbed into the unobservable bare parameter of the Lagrangian.

The bubble: coupling renormalization

The one-loop correction to $2 \rightarrow 2$ scattering involves bubble diagrams in the s , t , and u channels:



Reading the s -channel diagram with the Feynman rules: two vertices $((-i\mu^\varepsilon\lambda)^2)$, two internal propagators, a loop integration, and a symmetry factor $S = 2$ (the two internal lines can be interchanged). Setting $P = p_1 + p_2$:

$$i\mathcal{M}_s = \frac{(-i\mu^\varepsilon\lambda)^2}{2} \int \frac{d^d\ell}{(2\pi)^d} \frac{i}{\ell^2 - m^2} \cdot \frac{i}{(\ell - P)^2 - m^2}. \quad (6.13)$$

Power counting. For large $|\ell|$, the integrand behaves as $1/\ell^4$:

$$\int^\Lambda \frac{d^4\ell}{\ell^4} \sim \int^\Lambda \frac{\ell^3 d\ell}{\ell^4} \sim \ln \Lambda. \quad (6.14)$$

The integral is **logarithmically divergent**.

Evaluation in dimensional regularization. One introduces a single Feynman parameter x to combine the two denominators, shifts the loop momentum $\ell \rightarrow \ell + xP$, and then applies the master integral (6.3) with $n = 2$ and $\Delta = m^2 - x(1-x)s$. The result is

$$I(s) = \mu^\varepsilon \int \frac{d^d\ell}{(2\pi)^d} \frac{1}{(\ell^2 - m^2)[(\ell - P)^2 - m^2]} = \frac{i}{16\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) - F(s) \right], \quad (6.15)$$

where $s = P^2$ and $F(s)$ is the finite kinematic function (which depends on the renormalization scale μ)

$$F(s) = \int_0^1 dx \ln \frac{m^2 - x(1-x)s}{\mu^2}. \quad (6.16)$$

The full details of the Feynman-parameter technique and the evaluation of J_2 are given in Appendix G.

Carefully tracking the signs, $(-i\mu^\varepsilon\lambda)^2 = -\mu^{2\varepsilon}\lambda^2$ from the vertices and $(i)(i) = -1$ from the two propagators combine to give an overall $+\mu^{2\varepsilon}\lambda^2/2$, so that $i\mathcal{M}_s = (\mu^\varepsilon\lambda^2/2)I(s)$. Summing the three channels gives

$$i\mathcal{M}^{(1\text{-loop})} = +\frac{i\mu^\varepsilon\lambda^2}{32\pi^2} \left[\frac{6}{\varepsilon} + \text{finite}(s, t, u) \right], \quad (6.17)$$

so that, to order λ^2 , the full $2 \rightarrow 2$ amplitude reads

$$i\mathcal{M}(s, t, u) = -i\mu^\varepsilon\lambda + \frac{i\mu^\varepsilon\lambda^2}{32\pi^2} \left[\frac{6}{\varepsilon} + \text{finite}(s, t, u) \right] + O(\lambda^3). \quad (6.18)$$

Bare vs physical coupling. Exactly as for the mass, the parameter λ that appears in the Lagrangian (6.1) is the *bare* coupling, and it is *not* the quantity measured in a scattering experiment. The physical (renormalized) coupling λ_{phys} is *defined* as the value of the amplitude $-i\mathcal{M}$ at some reference kinematic point, say the symmetric point $s_0 = t_0 = u_0 = 4m_{\text{phys}}^2/3$:

$$\lambda_{\text{phys}} \equiv -\mathcal{M}(s_0, t_0, u_0). \quad (6.19)$$

Reading this off from (6.18) and using $\lambda_{\text{phys}} = -\mathcal{M} = \mu^\varepsilon\lambda - (\mu^\varepsilon\lambda^2/32\pi^2)[6/\varepsilon + \text{finite}]$, one finds schematically

$$\lambda_{\text{phys}} = \mu^\varepsilon\lambda - \frac{\mu^\varepsilon\lambda^2}{32\pi^2} \left[\frac{6}{\varepsilon} + \text{finite} \right] + O(\lambda^3), \quad (6.20)$$

Since at leading order $\mu^\varepsilon\lambda = \lambda_{\text{phys}} + O(\lambda_{\text{phys}}^2)$, this relation can be inverted iteratively. Writing $\text{finite}_0 \equiv \text{finite}(s_0, t_0, u_0)$ for the finite part evaluated at the reference kinematic point, one obtains

$$\mu^\varepsilon\lambda = \lambda_{\text{phys}} + \frac{\lambda_{\text{phys}}^2}{32\pi^2} \left[\frac{6}{\varepsilon} + \text{finite}_0 \right] + O(\lambda_{\text{phys}}^3). \quad (6.21)$$

In particular, the singular part of the bare coupling is

$$\mu^\varepsilon\lambda = \lambda_{\text{phys}} + \frac{3\lambda_{\text{phys}}^2}{16\pi^2} \frac{1}{\varepsilon} + \text{finite},$$

so that λ_{phys} is *finite* precisely if the bare λ is divergent, with its divergence chosen to cancel the $6/\varepsilon$ pole of the bubble. This is the essence of **coupling renormalization**: just as for the mass, the divergence sits in the unobservable bare parameter, not in any measured quantity.

At this stage we have only established the existence of such a redefinition; the systematic machinery that implements it—counterterms, renormalization constants Z_λ , and the freedom in the finite part (the choice of *scheme*)—is developed in the next subsection.

Counterterms and the renormalized Lagrangian

We have seen that loop corrections produce UV-divergent shifts to the mass and the coupling constant. The strategy of **renormalization** is to recognize that the parameters m and λ appearing in the original Lagrangian (6.1) are not the physically measured quantities—they are **bare** (unobservable) parameters. The divergences are not a sickness of the theory; they are simply telling us that the relation between bare and physical parameters involves infinite corrections.

To make this systematic, we introduce **renormalization constants** Z_ϕ , Z_m , Z_λ and write the Lagrangian as

$$\mathcal{L} = \frac{1}{2}Z_\phi \partial_\mu \phi \partial^\mu \phi - \frac{1}{2}Z_m m^2 \phi^2 - \mu^\varepsilon \frac{\lambda Z_\lambda}{4!} \phi^4, \quad (6.22)$$

where m and λ are now the *renormalized* (finite) parameters, whose precise values depend on the renormalization scheme and scale. The perturbative expansion is counted in powers of this renormalized coupling λ , not in powers of the poles in ε . Accordingly, each renormalization factor has the structure

$$Z = 1 + \sum_{n=1}^{\infty} \lambda^n \sum_{k=1}^n \frac{c_{nk}}{\varepsilon^k},$$

with numerical coefficients c_{nk} fixed order by order. Thus a term such as λ^2/ε is still an $O(\lambda^2)$ contribution: the pole is part of the regulator-dependent coefficient, not an extra power of the coupling. These divergent coefficients are chosen so that, at each order in λ , the poles from the Z -factors cancel the poles coming from loop integrals before the limit $\varepsilon \rightarrow 0$ is taken.

The idea is best understood by splitting the Lagrangian into two pieces: $\mathcal{L} = \mathcal{L}_{\text{ren}} + \mathcal{L}_{\text{ct}}$, where

$$\mathcal{L}_{\text{ren}} = \frac{1}{2} \partial_\mu \phi \partial^\mu \phi - \frac{1}{2} m^2 \phi^2 - \mu^\varepsilon \frac{\lambda}{4!} \phi^4 \quad (6.23)$$

has exactly the same form as the original Lagrangian (but now with *renormalized* parameters), and $\mathcal{L}_{\text{ct}}$ collects the **counterterms**—the leftover pieces. Let us derive $\mathcal{L}_{\text{ct}}$ explicitly.

Step 1: Splitting the Lagrangian. Start from the bare Lagrangian (6.22) and peel off the renormalized part (6.23). Writing $Z_\phi = 1 + \delta_\phi$ and $Z_m = 1 - \delta_m$, one finds

$$\frac{1}{2}Z_\phi (\partial\phi)^2 = \frac{1}{2}(\partial\phi)^2 + \frac{1}{2}\delta_\phi (\partial\phi)^2, \quad (6.24)$$

$$-\frac{1}{2}Z_m m^2 \phi^2 = -\frac{1}{2}m^2 \phi^2 + \frac{1}{2}\delta_m m^2 \phi^2, \quad (6.25)$$

where the second line uses $Z_m = 1 - \delta_m$, i.e. $-(1 - \delta_m)m^2 = -m^2 + \delta_m m^2$. Similarly for the coupling: defining $\delta_\lambda = \lambda(Z_\lambda - 1)$, the quartic term splits as

$$-\mu^\varepsilon \frac{\lambda Z_\lambda}{4!} \phi^4 = -\mu^\varepsilon \frac{\lambda}{4!} \phi^4 - \mu^\varepsilon \frac{\delta_\lambda}{4!} \phi^4.$$

Collecting all the leftover pieces:

$$\mathcal{L}_{\text{ct}} = \frac{1}{2} \delta_\phi \partial_\mu \phi \partial^\mu \phi + \frac{1}{2} \delta_m m^2 \phi^2 - \mu^\varepsilon \frac{\delta_\lambda}{4!} \phi^4, \quad (6.26)$$

with

$$\delta_\phi = Z_\phi - 1, \quad \delta_m = 1 - Z_m, \quad \delta_\lambda = \lambda(Z_\lambda - 1). \quad (6.27)$$

Notice the sign asymmetry: $\delta_\phi = Z_\phi - 1$ but $\delta_m = 1 - Z_m$. This traces back to the fact that the kinetic and mass terms have *opposite* signs in the Lagrangian ($+\frac{1}{2}(\partial\phi)^2$ vs $-\frac{1}{2}m^2\phi^2$), so peeling off $\mathcal{L}_{\text{ren}}$ produces different signs for the two residuals. The payoff is that $\delta_m > 0$ (the counterterm adds a positive correction to the propagator), which will make the cancellation of divergences transparent.

Step 2: Deriving the counterterm Feynman rules. Since $\mathcal{L}_{\text{ct}}$ has the same monomial structure as $\mathcal{L}_{\text{ren}}$, it generates new vertices in the Feynman rules. The recipe is the same as for any interaction term in the Lagrangian:

1. Identify the relevant monomial in the interaction Lagrangian. Recall that in the Dyson series (Chapter 3), the n th-order term contains

$$(-i)^n \int d^4x_1 \cdots d^4x_n \mathcal{H}_I(x_1) \cdots \mathcal{H}_I(x_n),$$

and the interaction Hamiltonian density is $\mathcal{H}_I = -\mathcal{L}_{\text{int}}$. The two minus signs ($-i$ and $-\mathcal{L}$) combine to give a factor of $+i$ for each vertex: a term $c\phi^n$ in $\mathcal{L}_{\text{int}}$ produces a vertex factor starting with $i \times c$.

2. Go to momentum space, replacing each $\phi(x)$ by its Fourier transform $\tilde{\phi}(p)$ and each derivative ∂_μ by $-ip_\mu$. The integral $\int d^4x$ over the vertex position produces a momentum-conserving $(2\pi)^4 \delta^{(4)}(\sum p_i)$.
3. Count the combinatorial factor: n identical fields can be contracted with n external legs in $n!$ ways; this cancels the $1/n!$ in front of the monomial (if present).

Let us work through each counterterm.

- **Mass counterterm:** The monomial is $+\frac{1}{2}\delta_m m^2 \phi^2$. This is a two-point vertex (two ϕ -legs meeting at a single point). Going to momentum space: the field

$\phi(x)$ is replaced by $\tilde{\phi}(p) = \int d^4x \phi(x) e^{ip \cdot x}$, and a product $\phi(x)^2$ at the *same* spacetime point becomes, after integrating over x :

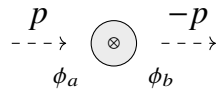
$$\int d^4x \phi(x)^2 = \int \frac{d^4p}{(2\pi)^4} \tilde{\phi}(p) \tilde{\phi}(-p).$$

(The δ -function from $\int d^4x e^{i(p_1+p_2) \cdot x}$ forces $p_2 = -p_1$.) Therefore the contribution to the action is:

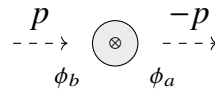
$$iS_{\text{ct}} \supset i \cdot \frac{1}{2} \delta_m m^2 \int \frac{d^4p}{(2\pi)^4} \tilde{\phi}(p) \tilde{\phi}(-p). \quad (6.28)$$

The vertex has two ϕ -fields (ϕ_a and ϕ_b , say) and two external legs carrying momenta p and $-p$. There are $2! = 2$ ways to assign which field contracts with which leg:

Assignment 1:



Assignment 2:



Both assignments give the same integral, so the Wick contractions produce a factor of 2. But the Lagrangian monomial carries $\frac{1}{2}$, so: $\frac{1}{2} \times 2 = 1$.

What about the fields $\tilde{\phi}(p)\tilde{\phi}(-p)$ and the integral $\int d^4p/(2\pi)^4$? These are *not* part of the vertex factor—they represent the external legs of the diagram (the incoming and outgoing propagators). The Feynman rule for a vertex is the “kernel” that remains after stripping off all external propagators and field operators. What is left is simply

$$\text{mass counterterm vertex:} \quad + i \delta_m m^2. \quad (6.29)$$

- **Field-strength counterterm:** The monomial is $+\frac{1}{2}\delta_\phi(\partial_\mu\phi)(\partial^\mu\phi)$. Going to momentum space, each derivative $\partial_\mu\phi(x)$ becomes $(-ip_\mu)\tilde{\phi}(p)$, so $(\partial_\mu\phi)(\partial^\mu\phi)$ produces a factor $(-ip_\mu)(ip^\mu) = p^2$, where p is the four-momentum flowing through the vertex. The same $\frac{1}{2}$ vs $2!$ cancellation as for the mass counterterm yields

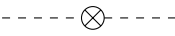
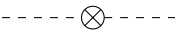
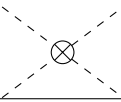
$$\text{field-strength counterterm vertex:} \quad + i \delta_\phi p^2. \quad (6.30)$$

Note that p^2 here is *not* set equal to m^2 : this counterterm is inserted on an internal propagator line where the particle is generically off-shell ($p^2 \neq m^2$). Its role is to correct the *momentum dependence* of the propagator (wave-function renormalization), unlike the mass counterterm which shifts only the constant part.

- **Coupling counterterm:** The monomial is $-\mu^\varepsilon \frac{\delta_\lambda}{4!} \phi^4$. This is a four-point vertex. Four external legs can connect to the four ϕ -fields in $4! = 24$ ways; the $\frac{1}{4!}$ cancels this, leaving

$$\text{coupling counterterm vertex:} \quad -i\mu^\varepsilon \delta_\lambda. \quad (6.31)$$

These **counterterm vertices** are drawn with a cross ($\otimes$) to distinguish them from ordinary vertices:

Counterterm	Diagram	Feynman rule
Mass (δ_m)		$i\delta_m m^2$
Field strength (δ_ϕ)		$i\delta_\phi p^2$
Coupling (δ_λ)		$-i\mu^\varepsilon \delta_\lambda$

Step 3: Consistency check—does the propagator come out right? The mass and field-strength counterterms both insert on the propagator line, giving a combined two-point vertex factor

$$V_{\text{ct}}^{(2)} = i(\delta_\phi p^2 + \delta_m m^2). \quad (6.32)$$

Inserting this into the Dyson series for the propagator:

$$G(p) = G_0 + G_0 \cdot V_{\text{ct}}^{(2)} \cdot G_0 + \cdots = \frac{i}{p^2 - m^2 + \delta_\phi p^2 + \delta_m m^2} = \frac{i}{Z_\phi p^2 - Z_m m^2}, \quad (6.33)$$

which is indeed the propagator from the bare Lagrangian (6.22). The two approaches—working directly with the bare Lagrangian or treating counterterms as extra Feynman rules—give the same result, as they must.

Step 4: How counterterms cancel divergences, order by order.

The key idea is simple: counterterm vertices are used *exactly like ordinary vertices*—insert them into diagrams, multiply by the appropriate factor, and integrate over internal momenta as usual. At each order in perturbation theory, we *choose* the counterterm coefficients (δ_m , δ_ϕ , δ_λ) so that the sum of all diagrams at that order is UV finite. Let us see how this works explicitly at order λ and order λ^2 .

At order λ : mass renormalization. The only one-loop correction to the propagator is the tadpole. It contributes a self-energy insertion $-i\Sigma_{\text{tadpole}}$, which is UV divergent (it contains a $2/\varepsilon$ pole). At the *same order* in λ , the mass counterterm

$$\text{renormalized: } \text{---}\text{---}\text{---}\overset{\circ}{\bullet}\text{---}\text{---}\text{---} + \text{---}\text{---}\text{---}\otimes\text{---}\text{---}\text{---} = \text{finite.} \quad (6.34)$$
$$-i\Sigma_{\text{tadpole}} + i\delta_m m^2 = \frac{-i\lambda m^2}{32\pi^2} \left[\frac{2}{\varepsilon} + \text{finite} \right] + i\delta_m m^2. \quad (6.35)$$

At order λ^2 : coupling renormalization. The one-loop correction to $2 \rightarrow 2$ scattering comes from the three bubble diagrams (the s -, t -, and u -channel diagrams derived above). Each bubble is $O(\lambda^2)$ and logarithmically divergent. At the same order, the coupling counterterm δ_λ (which is $O(\lambda^2)$) contributes a single tree-level diagram with a $\otimes$ vertex:

Equivalently, at order λ^2 one can write explicitly

so the UV pole in $i\mathcal{M}^{(1\text{-loop})}$ is cancelled by choosing δ_λ with the same coefficient and opposite sign. The three bubbles together give a divergence proportional to $3\lambda^2/(16\pi^2\varepsilon)$ (the factor of 3 because there are three channels). The counterterm vertex contributes $-i\mu^\varepsilon\delta_\lambda$. The divergence cancels if $\delta_\lambda = \frac{3\lambda^2}{32\pi^2} \left[\frac{2}{\varepsilon} + \text{finite} \right]$. The remaining finite part is a measurable correction to the scattering amplitude—the one-loop renormalization of the coupling constant.

Using the one-loop results from the tadpole (6.9) and the bubble (6.15), we can read off the renormalization constants at one loop:

$$Z_m = 1 - \frac{\lambda}{32\pi^2} \left(\frac{2}{\varepsilon} + \text{finite} \right) + O(\lambda^2), \quad (6.39)$$

$$Z_\lambda = 1 + \frac{3\lambda}{32\pi^2} \left(\frac{2}{\varepsilon} + \text{finite} \right) + O(\lambda^2). \quad (6.40)$$

Since $Z_m < 1$, we have $\delta_m = 1 - Z_m > 0$: the counterterm correction is positive, precisely compensating the positive loop self-energy. The fact that $Z_\phi = 1$ at one loop reflects the absence of a one-loop correction to the propagator's momentum dependence in ϕ^4 theory (the tadpole is momentum-independent). The “finite” parts depend on the choice of renormalization scheme, as we discuss next.

Renormalization schemes

The counterterms must cancel the $2/\varepsilon$ divergences—but nothing prevents us from also subtracting some finite amount. This freedom is not a bug; it simply reflects a choice of *what we mean* by m and λ . Different conventions define different **renormalization schemes**.

To see concretely how this works, consider the mass counterterm. From (6.9), the divergent self-energy is

$$\Sigma_{\text{tadpole}} = \frac{\lambda m^2}{32\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) - \ln \frac{m^2}{\mu^2} + 1 \right]. \quad (6.41)$$

We must choose δ_m to cancel the divergence. The simplest option: subtract *only* the pole.

- **MS (minimal subtraction):** Set

$$\delta_m^{\text{MS}} = \frac{\lambda}{32\pi^2} \frac{2}{\varepsilon}. \quad (6.42)$$

The counterterm removes the divergence and nothing else. The renormalized mass m then differs from the physical mass by a finite, calculable shift.

- **$\overline{\text{MS}}$ (modified minimal subtraction):** In practice, the combination $2/\varepsilon - \gamma_E + \ln(4\pi)$ appears in *every* dimensionally regularized integral. It is convenient to subtract all of it at once:

$$\delta_m^{\overline{\text{MS}}} = \frac{\lambda}{32\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) \right]. \quad (6.43)$$

This is by far the most widely used scheme in modern particle physics.

- **On-shell scheme:** Subtract *everything* so that the renormalized mass equals the physical mass (the pole of the propagator):

$$\delta_m^{\text{OS}} = \frac{\lambda}{32\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) - \ln \frac{m^2}{\mu^2} + 1 \right], \quad (6.44)$$

i.e. $\delta_m^{\text{OS}} = \Sigma_{\text{tadpole}}/m^2$, so that the correction to the propagator vanishes identically at $p^2 = m^2$.

All three schemes remove the divergence; they differ only in how much finite “leftover” is absorbed into the definition of m . But *physical observables cannot depend on this bookkeeping choice*. Let us prove this explicitly.

Explicit proof: the physical mass is scheme-independent.

The physical (pole) mass m_{phys}^2 is defined as the location of the pole of the full propagator:

$$m_{\text{phys}}^2 = m^2 + \Sigma_{\text{tadpole}} - \delta_m m^2, \quad (6.45)$$

where m is the renormalized mass parameter and δ_m is the counterterm—both scheme-dependent. To connect this with the original bare theory, recall from (6.22) that the mass term in the Lagrangian is

$$-\frac{1}{2}Z_m m^2 \phi^2.$$

This is just the original bare mass term written in renormalized variables, so it is natural to *define*

$$m_0^2 \equiv Z_m m^2 = (1 - \delta_m) m^2.$$

Thus m_0^2 is not a new physical quantity introduced out of nowhere: it is simply the **bare mass parameter** of the theory, rewritten in terms of the renormalized mass and the counterterm. The key point is that, for a fixed regularized theory, m_0^2 is held fixed while different renormalization schemes merely reshuffle how much of it is called “renormalized parameter” and how much is called “counterterm”:

$$m_0^2 = (1 - \delta_m) m^2 \implies m_{\text{phys}}^2 = m_0^2 + \Sigma_{\text{tadpole}}. \quad (6.46)$$

At the one-loop level, m_0 and Σ_{tadpole} (computed with the bare propagator) do not depend on the scheme, so neither does m_{phys} . (At higher orders, one must be more careful: the loop integrals themselves depend on which mass appears in the propagator, and the scheme independence of m_{phys} follows from the all-orders structure of the counterterms.) Let us verify scheme independence explicitly at one loop by direct calculation in two different schemes.

In the $\overline{\text{MS}}$ scheme: The renormalized mass is $m_{\overline{\text{MS}}}$ and

$$\delta_m^{\overline{\text{MS}}} = \frac{\lambda}{32\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) \right].$$

The renormalized self-energy (tadpole minus counterterm) is:

$$\begin{aligned} \Sigma_{\text{ren}}^{\overline{\text{MS}}} &= \Sigma_{\text{tadpole}} - \delta_m^{\overline{\text{MS}}} m_{\overline{\text{MS}}}^2 \\ &= \frac{\lambda m_{\overline{\text{MS}}}^2}{32\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) - \ln \frac{m_{\overline{\text{MS}}}^2}{\mu^2} + 1 \right] - \frac{\lambda m_{\overline{\text{MS}}}^2}{32\pi^2} \left[\frac{2}{\varepsilon} - \gamma_E + \ln(4\pi) \right] \\ &= \frac{\lambda m_{\overline{\text{MS}}}^2}{32\pi^2} \left[1 - \ln \frac{m_{\overline{\text{MS}}}^2}{\mu^2} \right]. \end{aligned} \quad (6.47)$$

The $1/\varepsilon$ pole has cancelled, and we are left with a *finite* correction. The physical mass is:

$$m_{\text{phys}}^2 = m_{\overline{\text{MS}}}^2 + \frac{\lambda m_{\overline{\text{MS}}}^2}{32\pi^2} \left[1 - \ln \frac{m_{\overline{\text{MS}}}^2}{\mu^2} \right] + O(\lambda^2). \quad (6.48)$$

In the on-shell scheme: By construction, $\delta_m^{\text{OS}} = \Sigma_{\text{tadpole}}/m_{\text{OS}}^2$, so:

$$\Sigma_{\text{ren}}^{\text{OS}} = \Sigma_{\text{tadpole}} - \delta_m^{\text{OS}} m_{\text{OS}}^2 = \Sigma_{\text{tadpole}} - \Sigma_{\text{tadpole}} = 0.$$

The physical mass is:

$$m_{\text{phys}}^2 = m_{\text{OS}}^2 + 0 = m_{\text{OS}}^2. \quad (6.49)$$

Are these the same? Yes! From (6.48), the $\overline{\text{MS}}$ mass is related to the on-shell mass by

$$m_{\text{OS}}^2 = m_{\overline{\text{MS}}}^2 + \frac{\lambda m_{\overline{\text{MS}}}^2}{32\pi^2} \left[1 - \ln \frac{m_{\overline{\text{MS}}}^2}{\mu^2} \right] + O(\lambda^2), \quad (6.50)$$

which is precisely (6.48). Both routes give the same m_{phys}^2 —they must, since m_{phys} is what an experiment measures.

The physical picture. A renormalization scheme is nothing but a **bookkeeping convention**: it dictates how to split a fixed, scheme-independent total into “parameter” and “loop correction.”

	Renormalized parameter	Loop correction
On-shell	$m_{\text{OS}}^2 = m_{\text{phys}}^2$	$\Sigma_{\text{ren}}^{\text{OS}} = 0$
$\overline{\text{MS}}$	$m_{\overline{\text{MS}}}^2 \neq m_{\text{phys}}^2$	$\Sigma_{\text{ren}}^{\overline{\text{MS}}} \neq 0$
Sum (physical)	$m_{\text{phys}}^2 = m^2 + \Sigma_{\text{ren}}$ (same in both cases)	

In the on-shell scheme, the parameter does all the work and the loop correction vanishes; in $\overline{\text{MS}}$, the parameter absorbs less and the loop correction carries the rest. The observable is always the *sum*, which is scheme-independent.

Box 6.1 **Summary: Renormalization of $\lambda\phi^4$**
Divergent diagrams:

- Tadpole (mass correction): quadratically divergent $\sim \Lambda^2$
- Bubble (vertex correction): logarithmically divergent $\sim \ln \Lambda$

Counterterms: Z_m cancels mass divergence, Z_λ cancels coupling divergence^a

Key lesson: The same structure applies to QED, but with the complication of spin and gauge invariance. The Ward identity constrains the counterterms, leaving only three independent renormalization constants.

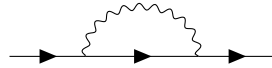
^aBecause the renormalization scale μ is arbitrary, the bare coupling λ_0 cannot depend on it. Requiring $\mu d\lambda_0/d\mu = 0$ implies that the renormalized coupling λ must *run* with μ ; the rate of change is the **beta function** $\beta(\lambda) = \mu d\lambda/d\mu = 3\lambda^2/(16\pi^2) + O(\lambda^3)$. Since $\beta > 0$, the coupling grows with energy and eventually hits a *Landau pole*, signalling the breakdown of perturbation theory: ϕ^4 theory is not asymptotically free. We will return to the beta function in detail when discussing QED.

Introduction: why QED needs renormalization

When we compute higher-order corrections in quantum electrodynamics, we encounter loop integrals that diverge in the ultraviolet (UV). These divergences are not a sign that the theory is sick; rather, they signal that the “bare” parameters in the Lagrangian (e_0, m_0)—that is, the charge and mass appearing in the unrenormalized Lagrangian—are not the physical quantities we measure. Renormalization is the systematic procedure that absorbs these infinities into redefinitions of the parameters and fields, leaving finite predictions for observables.

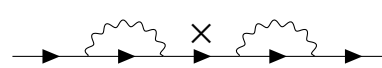
One-particle irreducible (1PI) diagrams. Recall that a Feynman diagram is called **1PI** if it *cannot* be disconnected by cutting a single internal line (we already encountered this concept when resumming the self-energy in (6.10)). Diagrams that *can* be split by cutting one line are called *one-particle reducible* (1PR).

1PI (irreducible)



Cannot be split by
cutting one line

1PR (reducible)



Cutting the middle line
disconnects the diagram

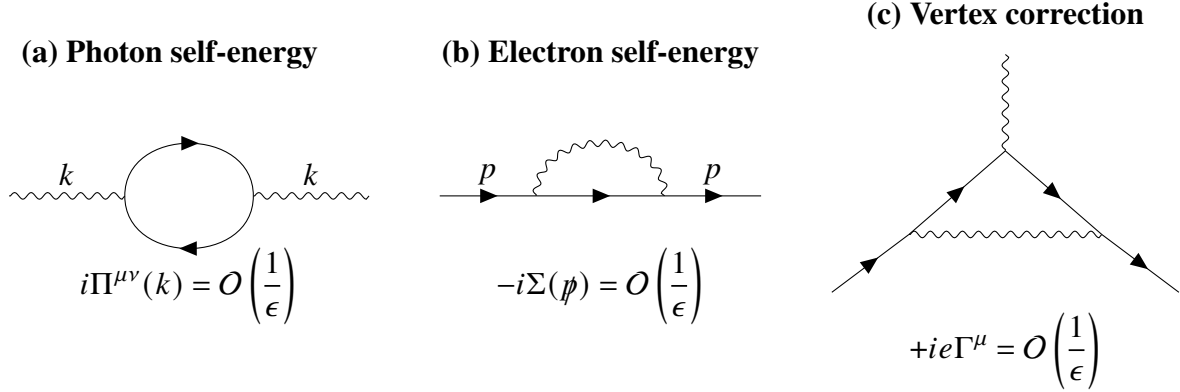
Why 1PI diagrams matter. The 1PI diagrams are the fundamental building blocks of perturbation theory for several reasons:

1. **No double counting:** Any Feynman diagram can be uniquely decomposed into 1PI subdiagrams connected by single propagators. Working with 1PI diagrams avoids overcounting contributions.
2. **Generating functionals:** The *effective action* $\Gamma[\phi]$ (the Legendre transform of the generating functional $W[J]$) generates precisely the 1PI diagrams. This is why Γ is also called the “1PI generating functional.”³
3. **Self-energies and vertex functions:** The physically meaningful corrections to propagators (self-energies Σ, Π) and vertices (Γ^μ) are defined as sums of 1PI diagrams. The full propagator is then obtained by resumming all 1PI insertions via a geometric series (Dyson equation).

³The idea, briefly, is the following. One introduces an external source $J(x)$ coupled linearly to the field, and defines a generating functional $Z[J]$ whose functional derivatives yield all correlation functions. Taking $W[J] = -i \ln Z[J]$ selects only the *connected* diagrams. One then defines the *classical field* $\phi_c(x) \equiv \delta W / \delta J(x)$, which represents the vacuum expectation value of the field in the presence of J . The *effective action* $\Gamma[\phi_c]$ is the Legendre transform of $W[J]$: $\Gamma[\phi_c] = W[J] - \int d^4x J(x) \phi_c(x)$. The key result is that the functional derivatives of Γ evaluated at $J = 0$ generate precisely the 1PI vertex functions—the amputated, connected Green functions that cannot be disconnected by cutting a single internal line. For a standard derivation, see Peskin–Schroeder [6], Ch. 11.

4. **Renormalization:** Only the 1PI Green functions need to be renormalized. Once the divergences in 1PI diagrams are cancelled by counterterms, all higher-order diagrams (built from 1PI blocks) become automatically finite.

The three primitively divergent one-loop diagrams in QED. At one loop, only three classes of 1PI diagrams are UV divergent by power counting. These are the building blocks from which all higher-order divergences are constructed:



A note on signs. The relative sign in the definitions $ie_0^2\Pi^{\mu\nu}(k)$ and $-i\Sigma(\not{p})$ is a matter of convention, chosen so that the Dyson-resummed propagators take the standard forms $D^{-1}(k) \sim k^2 - e_0^2\Pi_T(k^2)$ for the photon and $S^{-1}(\not{p}) \sim \not{p} - m_0 - \Sigma(\not{p})$ for the electron. It does *not* reflect a physical sign difference between the two effects. Part of the asymmetry is also notational: for the photon we factor out an explicit e_0^2 from the insertion, whereas for the electron self-energy the coupling dependence is included directly in $\Sigma(\not{p})$.

These diagrams represent:

- **(a) Vacuum polarization** $\Pi^{\mu\nu}(k)$: a virtual e^+e^- pair modifies the photon propagator. This leads to *charge renormalization* and a running coupling $\alpha(Q^2)$, where for spacelike exchange $Q^2 \equiv -k^2 > 0$.
- **(b) Electron self-energy** $\Sigma(\not{p})$: the electron emits and reabsorbs a virtual photon. This shifts the electron mass ($m_0 \rightarrow m$) and requires *wave-function renormalization* (Z_2).
- **(c) Vertex correction** Γ^μ : quantum fluctuations modify the electron-photon coupling. Combined with the self-energy, this determines Z_1 and ensures gauge invariance via the Ward identity $Z_1 = Z_2$.

Naive power counting from the integrands. We can determine the degree of divergence of each diagram by inspecting the large- k behaviour of the loop integral, exactly as we did for ϕ^4 theory.

- **Photon self-energy** $\Pi^{\mu\nu}(q)$: one loop with two fermion propagators ($\sim 1/k$ each at large k , since they are linear in $\not{k}$):

$$\Pi^{\mu\nu} \sim \int \frac{d^4k}{(2\pi)^4} \frac{\text{tr}[\gamma^\mu \not{k} \gamma^\nu \not{k}]}{k^2 \cdot k^2} \sim \int^\Lambda \frac{k^3 dk \cdot k^2}{k^4} \sim \Lambda^2 \quad (\text{quadratically divergent}).$$

(6.51)

- **Electron self-energy** $\Sigma(p)$: one fermion and one photon propagator in the loop ($\sim 1/k$ and $\sim 1/k^2$ at large k):

$$\Sigma \sim \int \frac{d^4k}{(2\pi)^4} \frac{\gamma^\mu \not{k} \gamma_\mu}{k^2 \cdot k^2} \sim \int^\Lambda \frac{k^3 dk \cdot k}{k^4} \sim \Lambda \quad (\text{linearly divergent}). \quad (6.52)$$

- **Vertex correction** $\Lambda^\mu(p', p)$: two fermion and one photon propagator ($\sim 1/k^5$ in total at large k):

$$\Lambda^\mu \sim \int \frac{d^4k}{(2\pi)^4} \frac{\gamma^\nu \not{k} \gamma^\mu \not{k} \gamma_\nu}{k^2 \cdot k^2 \cdot k^2} \sim \int^\Lambda \frac{k^3 dk \cdot k^2}{k^6} \sim \ln \Lambda \quad (\text{logarithmically divergent}). \quad (6.53)$$

However, gauge invariance (the Ward–Takahashi identity) reduces the actual degree of divergence: the photon self-energy becomes logarithmic rather than quadratic, and the electron self-energy becomes logarithmic rather than linear. All three divergences are ultimately *logarithmic*, i.e., poles $\sim 1/\varepsilon$ in dimensional regularization. A systematic classification via the *superficial degree of divergence* will be developed in Section 6.7.

The counterterm strategy. The bare Lagrangian is rewritten in terms of renormalized fields and parameters plus counterterms:

$$\mathcal{L}_0 = \mathcal{L}_{\text{ren}} + \delta\mathcal{L}, \quad (6.54)$$

where $\delta\mathcal{L}$ contains $(Z_2 - 1)$, δm , $(Z_3 - 1)$, and $(Z_1 - 1)$ contributions that generate counterterm vertices. These counterterms are chosen so that the sum of loop diagrams plus counterterm insertions is UV finite. The following sections develop this program in detail for each of the three primitively divergent structures.

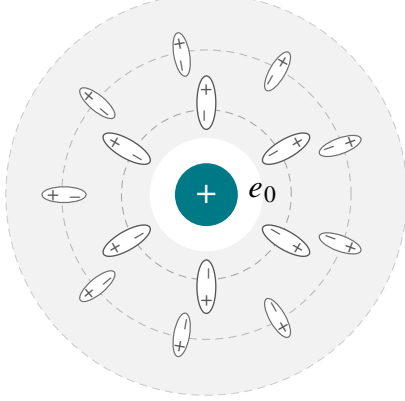
6.2 Photon self-energy in QED

Physical interpretation: vacuum polarization

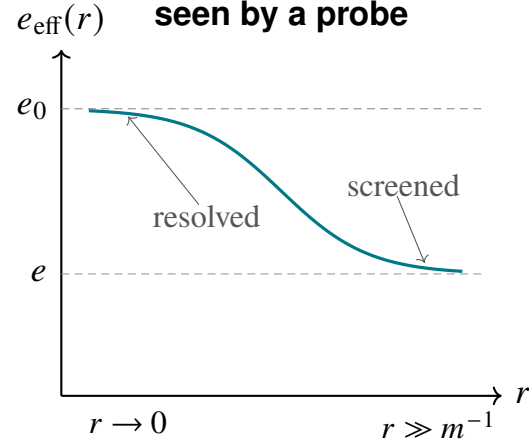
The photon self-energy, also known as *vacuum polarization*, represents one of the most profound effects in quantum electrodynamics: the quantum vacuum is not “empty” but rather a boiling sea of virtual electron-positron pairs that constantly pop in and out of existence. When a photon propagates through the vacuum, it can temporarily convert into such a virtual e^+e^- pair, which then annihilates back into a photon. This process modifies the effective photon propagator and has far-reaching physical consequences:

- **Charge screening:** At large distances (low momenta $|k^2| \ll m^2$), the virtual pairs partially screen the “bare” electric charge. The observed charge e is smaller than the bare charge e_0 that would be measured at infinitesimally short distances.

**Polarized vacuum
around a positive source**



**Effective charge
seen by a probe**



The left panel illustrates how virtual e^+e^- pairs orient themselves around a positive source charge: the negative ends of the dipoles are attracted inward, partially cancelling the field seen at large distances. The right panel shows the resulting effective charge $e_{\text{eff}}(r)$: at distances $r \gg 1/m$ (the electron Compton wavelength), the full screening is in effect and one measures $e_{\text{eff}} = e$; as $r \rightarrow 0$, the screening is penetrated and $e_{\text{eff}} \rightarrow e_0 > e$.

- **Running coupling:** At short distances (equivalently, for large spacelike momentum transfer $Q^2 \equiv -k^2 \gg m^2$), less screening occurs, so the effective coupling $\alpha(Q^2)$ *increases* with Q^2 . This is the famous “running of the QED coupling,” verified experimentally at LEP and other high-energy experiments.
- **Lamb shift contribution:** The vacuum polarization modifies the Coulomb potential at short distances, contributing (at order $\alpha^5 m$) to the Lamb shift in hydrogen. This was one of the first precision tests of QED.
- **Muon $g - 2$:** Vacuum polarization loops (including hadronic contributions) are crucial ingredients in the theoretical prediction of the anomalous magnetic moment of the muon.

The calculation below shows how the vacuum polarization tensor $\Pi^{\mu\nu}(k)$ arises from a one-loop Feynman diagram and how gauge invariance (the Ward identity) constrains it to be purely transverse: $\Pi^{\mu\nu}(k) = (k^2 g^{\mu\nu} - k^\mu k^\nu) \Pi(k^2)$.

From the Møller diagram to a self-energy insertion (explicit derivation and diagrams)

We illustrate how the photon self-energy (vacuum polarization) appears as an *insertion* into an internal photon line, starting from the tree-level (one-photon exchange) contribution to Møller scattering. For the purposes of this subsection, it

is enough to treat the external fermion lines as conserved currents; the same algebra applies to many other QED processes.

Tree-level amplitude as “current $\times$ propagator $\times$ current”.

Let p_1, p_2 be the incoming electron momenta and p_3, p_4 the outgoing ones. Define the momentum transfer

$$k \equiv p_1 - p_3 = p_4 - p_2. \quad (6.55)$$

The t -channel one-photon exchange amplitude can be written as

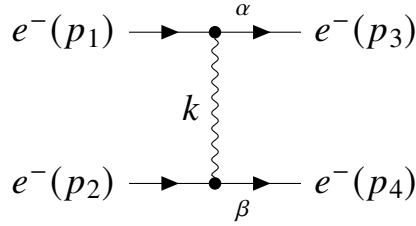
$$i\mathcal{M}^{(0)} = (-ie_0)^2 \left[\bar{u}(p_3) \gamma^\alpha u(p_1) \right] (iD_{F\alpha\beta}(k)) \left[\bar{u}(p_4) \gamma^\beta u(p_2) \right]. \quad (6.56)$$

It is convenient to introduce the external currents

$$J_1^\alpha(-k) \equiv \bar{u}(p_3) \gamma^\alpha u(p_1), \quad J_2^\beta(k) \equiv \bar{u}(p_4) \gamma^\beta u(p_2), \quad (6.57)$$

so that

$$i\mathcal{M}^{(0)} = (-ie_0)^2 J_1^\alpha(-k) (iD_{F\alpha\beta}(k)) J_2^\beta(k). \quad (6.58)$$



Tree-level t -channel exchange (one photon with momentum k).

The one-loop vacuum polarization insertion.

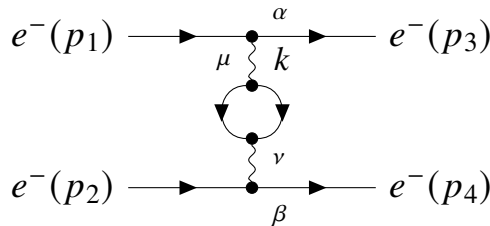
Now consider the correction where the internal photon propagator is modified by a vacuum polarization subdiagram. At order e_0^4 , the corresponding amplitude is

$$i\mathcal{M}_{\text{VP}}^{(1)} = (-ie_0)^2 J_1^\alpha(-k) (iD_{F\alpha\mu}(k)) (ie_0^2 \Pi^{\mu\nu}(k)) (iD_{F\nu\beta}(k)) J_2^\beta(k). \quad (6.59)$$

Comparing (6.58) and (6.59), we read off the *replacement rule* for an internal photon line (to this order):

$$iD_{F\alpha\beta}(k) \longrightarrow iD_{F\alpha\beta}(k) + iD_{F\alpha\mu}(k) (ie_0^2 \Pi^{\mu\nu}(k)) iD_{F\nu\beta}(k). \quad (6.60)$$

This is precisely the statement that the self-energy behaves like a 1PI *insertion* on the propagator.



Vacuum polarization insertion (fermion loop) on the exchanged photon.

Explicit expression for $\Pi^{\mu\nu}(k)$

The inserted blob is the 1PI photon two-point function. At one loop (electron loop), its definition in momentum space is

$$ie_0^2 \Pi^{\mu\nu}(k) = (-1) (-ie_0)^2 \int \frac{d^d p}{(2\pi)^d} \text{Tr} \left[\gamma^\mu \frac{i(\not{p} + m_0)}{p^2 - m_0^2 + i\epsilon} \gamma^\nu \frac{i(\not{p} + \not{k} + m_0)}{(p+k)^2 - m_0^2 + i\epsilon} \right]. \quad (6.61)$$

Collecting the factors of i and rearranging yields the more compact form

$$ie_0^2 \Pi^{\mu\nu}(k) = -e_0^2 \int \frac{d^d p}{(2\pi)^d} \frac{\text{Tr} \left[\gamma^\mu (\not{p} + m_0) \gamma^\nu (\not{p} + \not{k} + m_0) \right]}{(p^2 - m_0^2 + i\epsilon)((p+k)^2 - m_0^2 + i\epsilon)}. \quad (6.62)$$

The overall minus sign comes from the fermion loop.

Why the insertion modifies only the transverse propagator in physical amplitudes.

When the photon is exchanged between on-shell external fermions, the currents are conserved:

$$k_\alpha J_1^\alpha(-k) = 0, \quad k_\beta J_2^\beta(k) = 0, \quad (6.63)$$

so any part of the corrected propagator proportional to k_α or k_β drops out from the amplitude. In this context (physical amplitudes with conserved external currents), one can decompose $\Pi^{\mu\nu}(k)$ into transverse and longitudinal pieces and keep only the transverse part; the longitudinal part does not contribute. Note that this simplification applies specifically to amplitudes sandwiched between conserved currents; for off-shell Green functions or internal subdiagrams, the full tensor structure must be retained.

Geometric-series resummation of self-energy insertions

The vacuum-polarization subdiagram is a *1PI two-point insertion* on a photon line. Algebraically, it behaves like a matrix (in Lorentz indices) inserted between two free propagators. We write $iD_0^{\mu\nu}(k)$ for the free photon propagator and define the shorthand

$$(i\Pi)^{\mu\nu}(k) \equiv ie_0^2 \Pi^{\mu\nu}(k) \quad (6.64)$$

for the full 1PI vacuum-polarization insertion (i.e. the blob including the coupling factors e_0^2). The parentheses around $i\Pi$ distinguish this combined object from the stripped tensor $\Pi^{\mu\nu}$ (without e_0^2).

Step 1: write the series with repeated insertions. If we allow *any number* of vacuum-polarization insertions on the same internal photon line, the contribution to the effective propagator is the Neumann (geometric) series

$$iD^{\mu\nu}(k) = iD_0^{\mu\nu}(k) + iD_0^{\mu\alpha}(k) (i\Pi_{\alpha\beta}(k)) iD_0^{\beta\nu}(k) \\ + iD_0^{\mu\alpha}(k) (i\Pi_{\alpha\beta}(k)) iD_0^{\beta\gamma}(k) (i\Pi_{\gamma\delta}(k)) iD_0^{\delta\nu}(k) + \dots \quad (6.65)$$

Schematically, suppressing Lorentz indices:

$$D_0 + D_0(i\Pi)D_0 + D_0(i\Pi)D_0(i\Pi)D_0 + \dots$$

Step 2: factor out one free propagator and identify the geometric structure.

Define the Lorentz-index matrix (at fixed k)

$$(M(k))^\mu{}_\nu \equiv (i\Pi(k))^\mu{}_\alpha (iD_0(k))^\alpha{}_\nu. \quad (6.66)$$

Then each term in (6.65) has the form

$$iD_0(i\Pi iD_0)^n = iD_0 M^n, \quad (6.67)$$

so the full series becomes

$$iD = iD_0 \left(\mathbb{1} + M + M^2 + \dots \right). \quad (6.68)$$

Formally (as a matrix series), for sufficiently small coupling this sums to

$$\mathbb{1} + M + M^2 + \dots = (\mathbb{1} - M)^{-1}. \quad (6.69)$$

Substituting (6.69) into (6.68) gives

$$iD = iD_0 (\mathbb{1} - M)^{-1}. \quad (6.70)$$

Step 3: rewrite the result in the standard Dyson form. Using the definition (6.66),

$$\begin{aligned} \mathbb{1} - M &= \mathbb{1} - (i\Pi)(iD_0) = (iD_0)^{-1}(iD_0) - (i\Pi)(iD_0) \\ &= \left[(iD_0)^{-1} - (i\Pi) \right] (iD_0), \end{aligned} \quad (6.71)$$

so

$$(\mathbb{1} - M)^{-1} = (iD_0)^{-1} \left[(iD_0)^{-1} - (i\Pi) \right]^{-1}. \quad (6.72)$$

Plugging (6.72) into (6.70) yields the compact resummed propagator

$$iD^{\mu\nu}(k) = \left(\left[(iD_0)^{-1} - (i\Pi) \right]^{-1} \right)^{\mu\nu}(k) = \frac{iD_0^{\mu\alpha}(k)}{\delta_\alpha{}^\nu - (i\Pi)_{\alpha\beta}(k) iD_0^{\beta\nu}(k)}. \quad (6.73)$$

Here the final fraction is a compact notation for the inverse matrix in Lorentz-index space; explicitly,

$$\frac{1}{\delta_\alpha{}^\nu - (i\Pi)_{\alpha\beta} iD_0^{\beta\nu}} \equiv \left[\mathbb{1} - (i\Pi)(iD_0) \right]^{-1}{}_\alpha{}^\nu.$$

Equation (6.73) is the propagator-level Dyson equation. Recalling that $(i\Pi)^{\mu\nu} = ie_0^2\Pi^{\mu\nu}$, this reads

$$D^{-1} = D_0^{-1} - e_0^2 \Pi. \quad (6.74)$$

Step 4: check explicitly in Feynman gauge (scalar-like geometric sum).⁴ In Feynman gauge,

$$iD_0^{\mu\nu}(k) = \frac{-ig^{\mu\nu}}{k^2 + i\epsilon}, \quad (iD_0^{-1})_{\mu\nu}(k) = i(k^2 + i\epsilon)g_{\mu\nu}. \quad (6.75)$$

Why the geometric series becomes scalar: dropping longitudinal pieces.

Inside a physical scattering amplitude, the photon propagator is never used as a free-standing object: it is always *sandwiched between conserved currents*. This kills any longitudinal contribution proportional to the external momentum k^μ .

1. Decompose $\Pi^{\mu\nu}(k)$ into transverse and longitudinal parts. Lorentz covariance allows

$$\Pi^{\mu\nu}(k) = -g^{\mu\nu}A(k^2) + k^\mu k^\nu B(k^2). \quad (6.76)$$

Equivalently (and more transparently for gauge invariance), define the **transverse projector**

$$P_T^{\mu\nu}(k) \equiv g^{\mu\nu} - \frac{k^\mu k^\nu}{k^2}. \quad (6.77)$$

It is called “transverse” because it projects any vector onto the subspace orthogonal to k^μ : contracting with k_μ gives $k_\mu P_T^{\mu\nu} = k^\nu - k^\nu(k^2/k^2) = 0$. In other words, $P_T^{\mu\nu}$ kills the component of a vector along k and keeps only the part transverse to k . One easily verifies that it is indeed a projector: $P_T^{\mu\alpha} P_{T\alpha}{}^\nu = P_T^{\mu\nu}$.

Any symmetric tensor built from $g^{\mu\nu}$ and k^μ can be written as

$$\Pi^{\mu\nu}(k) = P_T^{\mu\nu}(k) \Pi_T(k^2) + \frac{k^\mu k^\nu}{k^2} \Pi_L(k^2), \quad (6.78)$$

for scalar functions Π_T, Π_L . The two parametrizations are related by simple linear combinations:

$$-g^{\mu\nu}A + k^\mu k^\nu B = \left(g^{\mu\nu} - \frac{k^\mu k^\nu}{k^2}\right) \Pi_T + \frac{k^\mu k^\nu}{k^2} \Pi_L. \quad (6.79)$$

⁴The free photon propagator in a general covariant gauge is $iD_0^{\mu\nu}(k) = \frac{-i}{k^2 + i\epsilon} \left[g^{\mu\nu} - (1 - \xi) \frac{k^\mu k^\nu}{k^2} \right]$, where ξ is the gauge parameter. **Feynman gauge** corresponds to $\xi = 1$, for which the propagator reduces to $-ig^{\mu\nu}/(k^2 + i\epsilon)$ —proportional to the metric, with no $k^\mu k^\nu$ term. This greatly simplifies calculations because the propagator acquires the same Lorentz structure as a scalar propagator (times $g^{\mu\nu}$), and the $k^\mu k^\nu$ terms that would generate complicated tensor integrals are absent. Physical results are gauge-independent (as guaranteed by the Ward identity), so we are free to choose the gauge that makes the calculation simplest.

In a gauge-invariant regularization (e.g. dimensional regularization) the Ward identity implies transversality of the vacuum polarization,

$$k_\mu \Pi^{\mu\nu}(k) = 0, \quad (6.80)$$

which forces $\Pi_L(k^2) = 0$, i.e. $\Pi^{\mu\nu}$ is purely transverse.

Box 6.2 Dictionary of notations for the vacuum polarization

Several equivalent parametrizations of the stripped tensor $\Pi^{\mu\nu}(k)$ are used in this section and in the literature. They are all related as follows.

(i) The A, B parametrization (Mandl–Shaw convention):

$$\Pi^{\mu\nu}(k) = -g^{\mu\nu} A(k^2) + k^\mu k^\nu B(k^2).$$

(ii) The transverse/longitudinal decomposition:

$$\Pi^{\mu\nu}(k) = P_T^{\mu\nu}(k) \Pi_T(k^2) + \frac{k^\mu k^\nu}{k^2} \Pi_L(k^2).$$

(iii) The gauge-invariant scalar function $\Pi(k^2)$, defined when $\Pi^{\mu\nu}$ is transverse:

$$\Pi^{\mu\nu}(k) = (k^2 g^{\mu\nu} - k^\mu k^\nu) \Pi(k^2).$$

Relations. Comparing (i) and (iii) term by term ($-g^{\mu\nu} A + k^\mu k^\nu B = k^2 g^{\mu\nu} \Pi - k^\mu k^\nu \Pi$):

$$A(k^2) = -k^2 \Pi(k^2), \quad B(k^2) = -\Pi(k^2) = \frac{A(k^2)}{k^2}.$$

And from (i) vs (ii): $\Pi_T = -A = k^2 \Pi$, $\Pi_L = -A + k^2 B = 0$ (when the Ward identity holds).

In the text below, we use primarily $A(k^2)$ (following Mandl–Shaw) and the relation $A(0) = 0$ (from gauge invariance). The key points are:

- In a current sandwich, only A matters (the B term drops out by current conservation).
- $A(k^2)$ vanishes at $k^2 = 0$ (photon stays massless).
- $A(k^2) = k^2 [A'(0) + \Pi_c(k^2)]$, where $A'(0)$ is the UV-divergent constant and Π_c the finite remainder.

2. Why longitudinal pieces drop out even without invoking (6.80). Take any amplitude in which a photon line of momentum k is exchanged between two

electromagnetic currents J_1^μ and J_2^ν . The relevant factor has the form

$$\mathcal{M} \supset J_{1\mu}(-k) D^{\mu\nu}(k) J_{2\nu}(k). \quad (6.81)$$

For on-shell external fermions, the currents are conserved:

$$k_\mu J_1^\mu(-k) = 0, \quad k_\nu J_2^\nu(k) = 0. \quad (6.82)$$

Indeed, using $k = p_1 - p_3 = p_4 - p_2$ and the on-shell Dirac equations $(\not{p} - m)u(p) = 0$ and $\bar{u}(p)(\not{p} - m) = 0$, we find

$$\begin{aligned} k_\mu J_1^\mu(-k) &= \bar{u}(p_3) \not{k} u(p_1) = \bar{u}(p_3) (\not{p}_1 - \not{p}_3) u(p_1) = \bar{u}(p_3) (m - m) u(p_1) = 0, \\ k_\nu J_2^\nu(k) &= \bar{u}(p_4) \not{k} u(p_2) = \bar{u}(p_4) (\not{p}_4 - \not{p}_2) u(p_2) = \bar{u}(p_4) (m - m) u(p_2) = 0. \end{aligned} \quad (6.83)$$

Now consider any term in $D^{\mu\nu}(k)$ proportional to k^μ or k^ν . For example, suppose

$$D^{\mu\nu}(k) = D_T(k) P_T^{\mu\nu}(k) + D_L(k) \frac{k^\mu k^\nu}{k^2}. \quad (6.84)$$

Then the longitudinal part contributes to (6.81) as

$$J_{1\mu}(-k) \left[D_L(k) \frac{k^\mu k^\nu}{k^2} \right] J_{2\nu}(k) = \frac{D_L(k)}{k^2} (J_1(-k) \cdot k) (k \cdot J_2(k)) = 0, \quad (6.85)$$

by (6.82). Therefore, *in any amplitude where the photon propagator is sandwiched between conserved currents*, longitudinal pieces do not contribute—independently of the gauge choice and even if the regularization produced an unphysical longitudinal component.

3. How this simplifies the geometric series in Feynman gauge. In Feynman gauge,

$$iD_0^{\mu\nu}(k) = \frac{-ig^{\mu\nu}}{k^2 + i\epsilon}. \quad (6.86)$$

When we insert $\Pi^{\mu\nu}(k)$ between propagators and then sandwich between conserved currents, we are free to replace $\Pi^{\mu\nu}(k)$ by its transverse part only:

$$\Pi^{\mu\nu}(k) \rightsquigarrow P_T^{\mu\nu}(k) \Pi_T(k^2), \quad (6.87)$$

because the difference is proportional to k^μ and/or k^ν and vanishes by (6.85). If one instead uses the A, B parametrization (6.76), then in a current sandwich

$$\begin{aligned} J_{1\mu} \Pi^{\mu\nu}(k) J_{2\nu} &= -A(k^2) J_1 \cdot J_2 + B(k^2) (J_1 \cdot k) (k \cdot J_2) \\ &= -A(k^2) J_1 \cdot J_2, \end{aligned} \quad (6.88)$$

so the B term drops out explicitly.

With (6.87), each insertion effectively carries the same Lorentz structure (transverse projector) and the series becomes “scalar-like”: the resummation acts only on the scalar coefficient of the transverse piece, while the longitudinal sector is irrelevant for the physical amplitude. Concretely, the dressed propagator in a conserved-current sandwich can be written as

$$iD^{\mu\nu}(k) \rightsquigarrow \frac{-i}{k^2 + i\epsilon - e_0^2 \Pi_T(k^2)} P_T^{\mu\nu}(k) + (\text{any longitudinal part}), \quad (6.89)$$

and the “any longitudinal part” does not contribute to (6.81) because of (6.85).

Bottom line. The phrase “drop longitudinal pieces” means: *for physical amplitudes built from conserved currents, terms proportional to k^μ or k^ν vanish upon contraction*. This makes the geometric-series resummation effectively scalar (it resums only the transverse scalar function), even though the full propagator is a tensor.

Then the n -th term in (6.65) reduces to a scalar geometric factor times $g^{\mu\nu}$:

$$iD_0 (i\Pi iD_0)^n \rightsquigarrow \frac{-ig^{\mu\nu}}{k^2 + i\epsilon} \left(\frac{-e_0^2 A(k^2)}{k^2 + i\epsilon} \right)^n. \quad (6.90)$$

To see this, note that each insertion gives the scalar factor $(i\Pi) \cdot (iD_0) \rightsquigarrow (-ie_0^2 A) \cdot \frac{-i}{k^2 + i\epsilon} = \frac{-e_0^2 A}{k^2 + i\epsilon}$. Summing $n = 0, 1, 2, \dots$ gives

$$\begin{aligned} iD^{\mu\nu}(k) &\rightsquigarrow \frac{-ig^{\mu\nu}}{k^2 + i\epsilon} \sum_{n=0}^{\infty} \left(\frac{-e_0^2 A(k^2)}{k^2 + i\epsilon} \right)^n \\ &= \frac{-ig^{\mu\nu}}{k^2 + i\epsilon} \frac{1}{1 + \frac{e_0^2 A(k^2)}{k^2 + i\epsilon}} = \frac{-ig^{\mu\nu}}{k^2 + i\epsilon + e_0^2 A(k^2)}. \end{aligned} \quad (6.91)$$

Expanding the final expression to first order in e_0^2 reproduces the first two terms $D_0 + D_0(i\Pi)D_0$ in (6.65), so the resummation is consistent order-by-order.

What this means physically. The geometric series is simply the statement that a propagator can contain *any number* of 1PI two-point insertions. Summing them reorganizes the perturbation expansion into a dressed propagator with an effective denominator, as in (6.73) (matrix form) and (6.91) (scalarized form in a conserved-current sandwich).

Regularization and gauge-invariant assumptions

The loop integral (6.62) defining the vacuum polarization is ultraviolet divergent for large loop momentum p . A simple illustrative regularization is to multiply the integrand by a convergence factor

$$f(p^2, \Lambda^2) = \left(\frac{-\Lambda^2}{p^2 - \Lambda^2} \right)^2, \quad (6.92)$$

where Λ is a cut-off parameter. For fixed large Λ , the integral becomes convergent; sending $\Lambda \rightarrow \infty$ formally restores the original theory.

In what follows we assume the theory has been regularized in a way that is compatible with gauge invariance and preserves a massless photon in the appropriate limit. In particular, we assume that after regularization all quantities we manipulate (notably the scalar functions introduced below) are finite, and that the photon remains massless (pole at $k^2 = 0$).

Explicit algebra for the modified propagator

Applying the results of the previous subsection, we now perform the explicit algebra in Feynman gauge. Using the replacement $\Pi^{\mu\nu}(k) \rightsquigarrow -g^{\mu\nu}A(k^2)$ (valid in physical amplitudes), the single-insertion formula (6.60) yields

$$\begin{aligned} \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} &\longrightarrow \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} + \frac{-ig_{\alpha\mu}}{k^2 + i\epsilon} \left(ie_0^2 [-g^{\mu\nu}A(k^2)] \right) \frac{-ig_{\nu\beta}}{k^2 + i\epsilon} \\ &= \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} + \frac{ie_0^2 A(k^2) g_{\alpha\beta}}{(k^2 + i\epsilon)^2} \\ &= \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} \left[1 - \frac{e_0^2 A(k^2)}{k^2 + i\epsilon} \right]. \end{aligned} \quad (6.93)$$

Since (after regularization) $A(k^2)$ is finite, the bracket in (6.93) can be recognized as the first terms in the expansion of a shifted denominator:

$$\begin{aligned} \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} \left[1 - \frac{e_0^2 A(k^2)}{k^2 + i\epsilon} \right] &= \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} \frac{k^2 + i\epsilon - e_0^2 A(k^2)}{k^2 + i\epsilon} \\ &= \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon + e_0^2 A(k^2)} + \mathcal{O}(e_0^4), \end{aligned} \quad (6.94)$$

which is the usual geometric-series resummation to this order.

Massless photon condition and the decomposition of $A(k^2)$

The bare photon propagator has a pole at $k^2 = 0$ (massless photon). After including vacuum polarization, the modified propagator (6.94) reads

$$\frac{-ig_{\alpha\beta}}{k^2 + i\epsilon + e_0^2 A(k^2)}. \quad (6.94)$$

This propagator has its pole where the denominator vanishes: $k^2 + e_0^2 A(k^2) = 0$. If the photon is to remain massless (pole at $k^2 = 0$), we need $e_0^2 A(0) = 0$, i.e.

$$A(0) = 0. \quad (6.95)$$

This condition is not an extra assumption: it follows from gauge invariance. In a gauge-invariant regularization (such as dimensional regularization), the Ward identity (6.80) forces the vacuum polarization to be purely transverse:

$$\Pi^{\mu\nu}(k) = (k^2 g^{\mu\nu} - k^\mu k^\nu) \Pi(k^2).$$

Comparing with the parametrization

$$\Pi^{\mu\nu}(k) = -g^{\mu\nu}A(k^2) + k^\mu k^\nu B(k^2),$$

we see that the coefficient of $g^{\mu\nu}$ must be proportional to k^2 :

$$A(k^2) = -k^2\Pi(k^2).$$

Hence $A(0) = 0$ automatically. Equivalently, if $A(0)$ were nonzero, the denominator of the propagator would behave near $k^2 = 0$ like $k^2 + e_0^2 A(0)$, which would shift the pole away from zero and act like a photon mass term. But a mass term $\frac{1}{2}m_\gamma^2 A_\mu A^\mu$ is not gauge invariant. Thus gauge invariance protects the photon mass: no photon mass term is generated by quantum corrections.

Therefore $A(k^2)$ must vanish at least linearly as $k^2 \rightarrow 0$, and we may write

$$A(k^2) = k^2 A'(0) + k^2 \Pi_c(k^2), \quad (6.96)$$

where

$$A'(0) \equiv A'(k^2 = 0) = \left. \frac{dA(k^2)}{dk^2} \right|_{k^2=0}, \quad (6.97)$$

and $\Pi_c(k^2)$ collects the terms beyond the linear one. In particular, $\Pi_c(0) = 0$, so that $k^2 \Pi_c(k^2)$ is of higher order in k^2 near the origin.

Charge renormalization extracted from the modified propagator

In scattering amplitudes, an internal photon propagator comes with two QED vertices, hence with a factor e_0^2 . For this reason it is convenient to consider the combination $e_0^2 D_F$, rather than the propagator by itself, and apply (6.93) to that combination. This is the object that actually appears in the lowest-order exchange amplitude, so the constant part of the vacuum-polarization correction can be absorbed into the definition of the renormalized electric charge. Multiplying (6.93) by e_0^2 and substituting (6.96) gives

$$\begin{aligned} \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} e_0^2 &\longrightarrow \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} e_0^2 \left[1 - \frac{e_0^2 A(k^2)}{k^2 + i\epsilon} \right] \\ &= \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} e_0^2 \left[1 - \frac{e_0^2 (k^2 A'(0) + k^2 \Pi_c(k^2))}{k^2 + i\epsilon} \right] \\ &= \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} e_0^2 \left[1 - e_0^2 A'(0) \right] + \frac{ig_{\alpha\beta}}{(k^2 + i\epsilon)^2} e_0^4 k^2 \Pi_c(k^2). \end{aligned} \quad (6.98)$$

The first term on the right-hand side is the original propagator multiplied by a constant factor. It is therefore natural to *define* the renormalized charge e by

$$e^2 \equiv e_0^2 \left[1 - e_0^2 A'(0) \right], \quad (6.99)$$

so that all lowest-order amplitudes expressed in terms of e automatically include the constant part of the photon self-energy correction.

Equivalently, introduce the renormalization constant Z_3 by

$$e^2 \equiv Z_3 e_0^2, \quad \text{or equivalently} \quad e \equiv Z_3^{1/2} e_0. \quad (6.100)$$

This is the four-dimensional form used in this subsection; in dimensional regularization one writes equivalently $\mu^{\epsilon/2} e = Z_3^{1/2} e_0$ once $Z_1 = Z_2$ is imposed. Comparing with (6.99), we read off

$$Z_3 = 1 - e_0^2 A'(0) + \mathcal{O}(e_0^4). \quad (6.101)$$

Expanding the square root then gives, to this order,

$$e = e_0 \sqrt{1 - e_0^2 A'(0)} = e_0 \left[1 - \frac{1}{2} e_0^2 A'(0) + \mathcal{O}(e_0^4) \right]. \quad (6.102)$$

This introduces no new physical content: it is simply the perturbative form of the definition of the renormalized charge. Its practical role is to show that, at one-loop order, bare and renormalized couplings differ only beyond the order to which we are working, so e_0 can be replaced by e consistently in higher-order terms.

We now want to express the result entirely in terms of the *renormalized* charge e rather than the bare charge e_0 . From (6.99), $e_0^2 = e^2 + \mathcal{O}(e^4)$, i.e. e_0 and e differ only at order e^2 .

The first term of (6.98) is $e_0^2 [1 - e_0^2 A'(0)] / (k^2 + i\epsilon)$, which is $e^2 / (k^2 + i\epsilon)$ by definition (6.99)—here the replacement is exact.

The second term is already of order e_0^4 (it carries an explicit factor e_0^4). Replacing $e_0^4 \rightarrow e^4$ introduces an error of order $e_0^4 \times \mathcal{O}(e^2) = \mathcal{O}(e^6)$, which is beyond the accuracy of our one-loop calculation (order e^4). Therefore, to order e^4 , we may freely write $e_0^4 \rightarrow e^4$ in the second term.

The final result for the photon propagator (times the two charges) becomes

$$\frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} e_0^2 \longrightarrow \frac{-ig_{\alpha\beta}}{k^2 + i\epsilon} e^2 + \frac{ig_{\alpha\beta}}{(k^2 + i\epsilon)^2} e^4 k^2 \Pi_c(k^2) + \mathcal{O}(e^6). \quad (6.103)$$

Interpretation of the result.

Equation (6.103) separates the photon self-energy effect into:

- a *constant* renormalization of the charge, encoded in e^2 via (6.99) (this part contains the UV divergence through $A'(0)$);
- a momentum-dependent remainder proportional to $\Pi_c(k^2)$; after the renormalization implied above, this is the non-constant part of the vacuum polarization that survives in physical radiative corrections.

The explicit evaluation of $\Pi_c(k^2)$ (and of the divergent constant $A'(0)$ in dimensional regularization) is carried out in detail in Section 7.1.